

Dezentrale Sendeleistungsregelung zur Kapazitätssteigerung drahtloser Netze mit gemeinsam genutztem Übertragungskanal

Dissertation

zur Erlangung des Doktorgrads (Dr. rer. nat.)

der Mathematisch-Naturwissenschaftlichen Fakultät

der Rheinischen Friedrich-Wilhelms-Universität Bonn

vorgelegt von

Christian de Waal

aus Düsseldorf

Bonn, den 28. März 2006

Angefertigt mit Genehmigung der Mathematisch-Naturwissenschaftlichen Fakultät der Rheinischen Friedrich-Wilhelms-Universität Bonn.

Erstgutachter: Prof. Dr. Peter Martini, Rheinische Friedrich-Wilhelms-Universität Bonn
Zweitgutachter: Prof. Dr. Martin Mauve, Heinrich-Heine-Universität Düsseldorf

Datum der mündlichen Prüfung: 14. Juli 2006

Diese Dissertation ist auf dem Hochschulschriftenserver der ULB Bonn
http://hss.ulb.uni-bonn.de/diss_online
elektronisch publiziert.

Erscheinungsjahr: 2006

Danksagung

Diese Dissertation ist im Rahmen meiner Tätigkeit als wissenschaftlicher Mitarbeiter in der Arbeitsgruppe von Prof. Dr. Martini am Institut für Informatik IV der Rheinischen Friedrich-Wilhelms-Universität Bonn entstanden.

Mein größter Dank richtet sich an Prof. Dr. Martini, der mir die Promotion in seiner Arbeitsgruppe ermöglicht hat. Seine zielgerichteten Fragen und Kommentare in Vorträgen und Besprechungen haben einen wesentlichen Beitrag zur Qualität der vorliegenden Dissertation geleistet.

Ich bin außerdem sehr darüber erfreut, dass ich mit Prof. Dr. Mauve einen Zweitgutachter gewinnen konnte, der schon lange erfolgreich in dem Themengebiet der drahtlosen Ad-hoc-Netze forscht. Ich möchte ihm ganz herzlich dafür danken, dass er sich die Zeit genommen hat, um sich mit meiner Dissertation auseinanderzusetzen.

Weiterhin danke ich Dr. Paul James und dem Nokia Research Center Bochum, an dem er seinerzeit tätig war, für die Möglichkeit, in dem BMBF-geförderten Forschungsprojekt „IPonAir“ mitarbeiten zu können. Im Kontext dieses Projekts sind viele der Ideen entstanden, die in dieser Dissertation niedergeschrieben sind.

Einen herzlichen Dank möchte ich meinen vielen jetzigen und ehemaligen wissenschaftlichen Mitarbeiter-Kollegen Dr. Markus Albrecht, Nils Aschenbruck, Dr. Simon Baatz, Christoph Barz, Dr. Matthias Frank, Elmar Gerhards-Padilla, Dr. Michael Gerharz, Rolf Göppfarth, Dr. Wolfgang Hansmann, Alexander Hertwig, Michael Köster, Felix Leder, Sascha Lettgen, Wolfgang Moll, Patrick Peschlow, Markus Pilz, Marc Plaggemeier, Lukas Pustina, Christoph Scholz, Simon Schwarzer, Dr. Jens Tölle und Dr. Andre Wenzel für ihre Unterstützung in fachlichen Fragen, für die generelle Unterstützung in anderen Belangen der täglichen Arbeit und nicht zuletzt für die netten Kaffeerunden aussprechen. Zwei meiner Kollegen sollen an dieser Stelle besonders hervorgehoben werden: Michael Gerharz danke ich auch für die effektive Zusammenarbeit im Themengebiet der drahtlosen Ad-hoc-Netze, von der wir sicherlich beide sehr profitiert haben. Patrick Peschlow danke ich dafür, dass er sich zu einer Diplomarbeit unter meiner Betreuung entschieden hat, in deren Rahmen er auch seinen Teil zu meiner Dissertation beigetragen hat.

Auch den weiteren Mitarbeitern Günter Feldt, Udo Fink, Elisabeth Kirsch, Julia Kouchaki, Christian Kühn, Erika Müller-Hilckmann und Karlheinz Schumacher danke ich für ihren Beitrag zur Atmosphäre und zum Funktionieren unserer Arbeitsgruppe.

Danksagung

Meinem Freund Nael Al-Anaswah danke ich für die Durchsicht der Dissertation und seine hilfreichen Anmerkungen.

Einen großen Dank möchte ich schließlich auch meiner Frau Bettina für ihre Geduld bei der Suche nach Fehlern, für ihre Geduld mit meiner Person und den Rückhalt, den sie mir dadurch gegeben hat, aussprechen.

Bonn, im Juli 2006

Christian de Waal

Zusammenfassung

Die Topologie eines drahtlosen Netzes wird maßgeblich durch die Sendeleistungen der beteiligten Geräte beeinflusst. Sind die Sendeleistungen zu gering, so zerfällt die Topologie in mehrere Zusammenhangskomponenten, zwischen denen keine Kommunikation möglich ist. Zu hohe Sendeleistungen haben jedoch den Nachteil, dass unnötig viele Geräte um das gemeinsam genutzte Übertragungsmedium konkurrieren. Daher kann es sich vorteilhaft auf die Netzkapazität auswirken, die Sendeleistungen der Geräte dynamisch anzupassen, so dass die Konnektivität gewahrt bleibt, während eine möglichst hohe Anzahl zeitlich überlappend stattfindender Übertragungen zugelassen wird.

Die Dissertation leistet zwei Beiträge zu diesem Themengebiet. Zunächst wird der Zusammenhang zwischen der Topologie und der Kapazität eines Netzes untersucht. Dann wird ein Verfahren zur dezentralen Sendeleistungsregelung entwickelt, das die Netztopologie einer bestimmten Graphklasse annähert, von der zuvor gezeigt wurde, dass sie ein hohes Kapazitätspotenzial aufweist.

Für die Untersuchung von Netztopologien kommt ein neuer Graphalgorithmus zur Approximation der Transportkapazität zum Einsatz, der in der Dissertation vorgestellt wird. Es handelt sich dabei um den ersten Algorithmus, der den möglichen Durchsatz einer beliebigen Menge von Datenströmen unter Einhaltung des Max-Min-Fairness-Kriteriums auch in großen Topologien mit mehreren tausend Knoten in akzeptabler Rechenzeit approximieren kann.

Anhand von Topologien, die dadurch entstehen, dass allen Stationen die gleiche Sendereichweite zugeordnet wird, werden einige grundsätzlichen Erkenntnisse zum Einfluss der gewählten Sendeleistungen auf die Netzkapazität erarbeitet. Dann wird die Kapazität zweier Topologieklassen untersucht, die sich aus praktischer Sicht gut zur verteilten Sendeleistungsregelung eignen. Dabei zeigt sich, dass die *Nearest-Neighbours*- den *Max-Degree-Topologien* vorzuziehen sind. Schließlich wird festgestellt, dass einige andere Topologieklassen, die in der Literatur zu Sendeleistungsregelung in drahtlosen Netzen zu finden sind, hinsichtlich ihrer Kapazität keine Vorteile gegenüber *Nearest-Neighbours-Topologien* bringen.

Da bisher kein praktikables verteiltes *Nearest-Neighbours*-Verfahren zur Sendeleistungsregelung in drahtlosen Netzen veröffentlicht wurde, wird der im Rahmen der Dissertation entwickelte *Cooperative Nearest-Neighbours Topology Control Algorithm* in verschiedenen Varianten vorgestellt.

Anhand von Simulationsergebnissen wird schließlich gezeigt, dass dieser Algorithmus auch

unter realistischen Bedingungen in vielen Fällen zu einer hohen Netzkapazität führt. Dabei werden auch die Stärken und Schwächen der unterschiedlichen Varianten des vorgestellten Verfahrens unter verschiedenen Bedingungen genau beleuchtet. Da eine dieser Varianten ein Max-Degree-Verfahren repräsentiert, kann gezeigt werden, dass Nearest-Neighbours-Topologien auch aus praktischer Sicht den Max-Degree-Topologien vorzuziehen sind.

Inhaltsverzeichnis

1	Einleitung	1
2	Grundlagen	5
2.1	Drahtlose Datenübertragung	5
2.1.1	Antennen	5
2.1.2	Bedingungen für eine erfolgreiche Datenübertragung	6
2.1.3	Die Ausbreitung von Funkwellen	6
2.1.3.1	Large-scale-Fading	7
2.1.3.2	Small-scale-Fading	8
2.2	Die Verbindungsschicht in drahtlosen Netzen	10
2.2.1	IEEE 802.11	11
2.2.2	Drahtlose Links und Interferenz	14
2.3	Routing in drahtlosen Multihop-Netzen	16
2.3.1	Netzweites Fluten	17
2.3.2	Proaktives und reaktives Routing	17
2.3.3	Link-State- und Distance-Vector-Protokolle	18
2.3.4	Das AODV-Protokoll	19
2.4	Topologien und deren Eigenschaften	21
2.4.1	Modellierung von Stationen, Links und Interferenz	21
2.4.2	Der Dehnungsfaktor	23
2.4.3	Übertragungs-Scheduling	23
2.4.4	Kapazität	28
2.5	Modellierung von Mobilität	29
2.5.1	Das Random-Waypoint-Modell	29
2.5.2	Das Gauß-Markov-Mobilitätsmodell	30
2.6	Fairness	33
3	Annahmen, Ziele und Anforderungen	37
3.1	Beschreibung der Anwendungsszenarien	37
3.2	Zusammenhang zwischen Sendeleistungsregelung und Topologieformung	37
3.3	Anforderungen an Verfahren zur Sendeleistungsregelung	39

4	Verwandte Arbeiten	43
4.1	Auswirkungen der Sendeleistungsregelung	43
4.1.1	Analytische Arbeiten	43
4.1.2	Eigenschaften konkreter Topologien	45
4.1.2.1	Broadcast-Kapazität	45
4.1.2.2	Transportkapazität	46
4.1.3	Zusammenfassung und Motivation neuer Verfahren zur Kapazitätsabschätzung	47
4.2	Ansätze zur Topologieerzeugung	48
4.2.1	Nur auf einem Distanzmaß basierende Ansätze	50
4.2.1.1	Max-Degree-Topologien	50
4.2.1.2	Nearest-Neighbours-Topologien	52
4.2.1.3	Minimale Spannbäume	55
4.2.1.4	Relative-Neighbourhood-Topologien	55
4.2.1.5	Gabriel-Graphen	56
4.2.1.6	Topologien mit vorgegebenem Dehnungsfaktor	57
4.2.1.7	Energie-effiziente Topologien	58
4.2.2	Geometrische Ansätze	58
4.2.2.1	Delaunay-Triangulation	59
4.2.2.2	Cone-based Topology Control	60
4.2.3	Andere Ansätze	61
4.2.3.1	COMPOW und ClusterPow	61
4.2.3.2	Berücksichtigung der aktuellen Netzlast	62
4.2.4	Zusammenfassung und Motivation des neuen Verfahrens	62
5	Verfahren zur Kapazitätsabschätzung von Topologien	65
5.1	Netzmodell	65
5.2	Abschätzung der Broadcast-Kapazität	66
5.3	Der Zwei-Phasen-Färbealgorithmus zur Abschätzung der Transportkapazität	67
5.3.1	Vorgehensweise des Algorithmus	68
5.3.2	Phase 1: Ressourcenverteilung unter Einhaltung strikter Fairness	69
5.3.3	Phase 2: Water-Filling	71
5.3.4	Färbeheuristik	72
5.3.5	Vergleich mit anderen Verfahren	73
6	Auswirkungen der Sendeleistungsregelung auf die Netzkapazität	75
6.1	Modellierung	75
6.1.1	Links	75
6.1.2	Interferenz	76
6.1.3	Verkehrsmuster	77
6.2	Untersuchte Strategien zur Topologieformung	78
6.3	Verwendete Metriken	82
6.3.1	Konnektivität	82

6.3.2	Routenoverhead	82
6.3.3	Inhomogenität von Sendeleistungszuweisungen	84
6.4	Broadcast-Kapazität	85
6.5	Transportkapazität	86
6.5.1	Common-Range-Topologien	87
6.5.2	NNTC- und Max-Degree-Topologien	91
6.5.3	Andere Topologien	93
6.6	Fazit	93
7	Ein verteiltes Verfahren zur Sendeleistungsregelung	95
7.1	Motivation	95
7.2	Arbeitsweise des Algorithmus	96
7.3	Sammeln lokaler Topologieinformation	98
7.4	Erhöhen der Sendeleistung	100
7.5	Verringern der Sendeleistung	101
7.5.1	Die einfache Variante	102
7.5.2	Signalisierung redundanter Links	102
7.5.3	Versuchsweises Reduzieren	103
7.6	Der Algorithmus im Überblick	104
8	Untersuchung der Mechanismen zur Linkerkennung und Sendeleistungsregelung	111
8.1	Beschreibung der Simulationsumgebung und der Szenarien	111
8.1.1	Bitübertragungsschicht	112
8.1.2	Verbindungsschicht	113
8.1.3	Vermittlungsschicht	114
8.1.4	Weitere Eigenschaften der beteiligten Geräte	114
8.1.5	Verkehrsmuster	115
8.2	Metriken	115
8.2.1	Abstand zwischen Zuweisungen von Sendeleistungen	115
8.2.2	Einschwingzeiten	116
8.2.3	Linkpersistenz	116
8.2.4	Kapazität	117
8.3	Untersuchung verschiedener Parameterbelegungen	119
8.3.1	Szenarien ohne Sendeleistungsregelung	119
8.3.1.1	Linkerkennung	119
8.3.1.2	Gerätedichte	125
8.3.1.3	Dynamik	126
8.3.2	Szenarien mit Sendeleistungsregelung	128
8.3.2.1	Zielintervall der Nachbarzahlen	128
8.3.2.2	Gerätedichte	128
8.3.2.3	Dynamik	129
8.4	Vergleich verschiedener Varianten der verteilten Sendeleistungsregelung . . .	130

8.4.1	Statische Szenarien ohne Smallscale-Fading	130
8.4.1.1	Auswirkung der Kooperation	131
8.4.1.2	Vergleich der Mechanismen zur Sendeleistungsreduktion .	134
8.4.2	Statische Szenarien mit Smallscale-Fading	135
8.4.2.1	Einschwingzeiten	135
8.4.2.2	Auswirkung der Kooperation	136
8.4.2.3	Vergleich der Mechanismen zur Sendeleistungsreduktion .	138
8.4.3	Dynamische Szenarien ohne Smallscale-Fading	139
8.4.3.1	Vergleich der Mechanismen zur Sendeleistungsreduktion .	139
8.4.3.2	Auswirkung der Kooperation	141
8.4.4	Dynamische Szenarien mit Smallscale-Fading	142
8.5	Fazit	142
9	Zusammenfassung und Ausblick	145
9.1	Zusammenfassung	145
9.2	Ausblick	147

Abbildungsverzeichnis

2.1	Veranschaulichung von durch Largescale- und Smallscale-Fading verursachter Dämpfung	6
2.2	Veranschaulichung der Konsequenzen von Multipfad-Ausbreitung	8
2.3	Signalstärkeschwankungsverteilungen für verschiedene Rice-Faktoren	9
2.4	Das Hidden-Node-Problem	13
2.5	Beispiel eines drahtlosen Multihop-Netzes	17
2.6	Veranschaulichung der Etablierung einer Route mit AODV	20
2.7	Veranschaulichung des Dehnungs-Begriffs	22
2.8	Beispiel eines sich periodisch wiederholenden TDMA-Schemas	24
2.9	Veranschaulichung verschiedener Graphfärbungen	26
2.10	Aufenthaltswahrscheinlichkeiten im eindimensionalen Random-Waypoint-Modell (ohne Verweilen am Ziel)	30
2.11	Einschränkung der Gerätepositionen auf vorgegebenes Rechteck	32
2.12	Beispiele für Bewegungen gemäß dem Gauß-Markov-Mobilitätsmodell	33
2.13	Beispiel zur Veranschaulichung von Fairness-Definitionen	33
3.1	Einfluss der Sendeleistungsregelung auf verschiedenen Ebenen des OSI-Referenzmodells	38
3.2	Verschärfung des Hidden-Node-Problems durch stark verschiedene Sendeleistungen in räumlicher Nähe zueinander	40
4.1	Sendeleistungsregelung als Abwägung zwischen entgegengesetzten Zielen	48
4.2	Schwäche des Max-Degree-Ansatzes	51
4.3	Die 4-Nearest-Neighbours-Mengen zweier Knoten	53
4.4	Schwäche des Nearest-Neighbours-Ansatzes	55
4.5	Bereiche um eine Kante, die in Relative-Neighbourhood- und Gabriel-Graphen keine Knoten enthalten	55
4.6	Beispiel-Triangulationen einer Menge von vier Punkten	59
4.7	Veranschaulichung des bei CBTC durch einen Nachbarn abgedeckten Bereichs	60
5.1	Beispiel für einen Topologiegraphen und Datenströme und den daraus abgeleiteten Konfliktgraphen	68
5.2	Der Zwei-Phasen-Färbealgorithmus	70

6.1	Visualisierung von Beispieltopologien für unterschiedliche Formungsstrategien	79
6.2	Zuordnung von Sendereichweiten in Max-Degree-Topologien	80
6.3	Broadcast-Kapazität mit verschiedenen Strategien zur Sendeleistungsregelung	85
6.4	Zweiteilung eines langen Datenstroms	86
6.5	Transportkapazität in Common-Range-Topologien mit 1000 Stationen	87
6.6	Einfluss verschiedener Größen auf die Transportkapazität (Skizze)	88
6.7	Transportkapazität in Szenarien mit 100 Stationen	89
6.8	Maximal erzielbarer Durchsatz in Common-Range-Topologien in Abhängig- keit von der Stationszahl	89
6.9	Spatial-Reuse und Routenoverhead in Common-Range-Topologien mit 1000 Stationen	90
6.10	Transportkapazität von NNTC- und Max-Degree-Topologien mit 1000 Stationen	91
6.11	Inhomogenität der Sendereichweiten in NNTC- und Max-Degree-Topologien mit 1000 Stationen	92
6.12	Transportkapazität in anderen Topologien	93
7.1	Zwei verschiedene Mengen redundanter Links eines Geräts	101
7.2	Prozedur <code>recv_beacon</code>	106
7.3	Prozedur <code>send_beacon</code>	107
7.4	Prozedur <code>increase_k</code>	107
7.5	Prozedur <code>decrease_k</code>	108
7.6	Prozedur <code>initialise_global_variables</code>	108
7.7	Grundgerüst des verteilten Algorithmus zur Sendeleistungsregelung	109
8.1	Anzahl bidirektionaler Links und Paketverluste in einem Beispielszenario im zeitlichen Verlauf	118
8.2	Erfolg von Datenübertragungen ohne und mit Mechanismus zur Linkerkennung	121
8.3	Verzögerung von Datenpaketen in erfolgreichen Strömen	122
8.4	Erfolg von Datenübertragungen für andere Linkerkennungsparameter	123
8.5	Erfolg von Datenübertragungen für unterschiedliche l_{win}	123
8.6	Linkpersistenz in Abhängigkeit von der Distanz	123
8.7	Linkpersistenz	124
8.8	Erfolg von Datenübertragungen in Szenarien mit anderer Gerätedichte	125
8.9	Erfolg von Datenübertragungen in dynamischen Szenarien	126
8.10	Erfolg von Datenübertragungen mit verteilter Sendeleistungsregelung	127
8.11	Einschwingzeiten für verschiedene Varianten des Algorithmus	131
8.12	Nachbarzahlen und Sendeleistungen in statischen Szenarien ohne Smallscale- Fading	132
8.13	Nachbarzahlen in statischen Szenarien mit geringer Gerätedichte ohne Smallscale-Fading	133
8.14	Sendeleistungen in statischen Szenarien mit geringer Gerätedichte ohne Smallscale-Fading	133
8.15	Sendeleistungen ohne versuchsweises Reduzieren in statischen Szenarien . . .	135

8.16	Verlauf der Topologieentwicklung in Szenarien mit Smallscale-Fading	136
8.17	Knotengrade in statischen Szenarien mit Smallscale-Fading	137
8.18	Knotengrade in dynamischen Szenarien ohne Smallscale-Fading	140
8.19	Sendeleistungen in dynamischen Szenarien ohne Smallscale-Fading	141

Tabellenverzeichnis

4.1	Übersicht über Verfahren zur Sendeleistungsregelung	63
5.1	Beispielfärbung des Konfliktgraphen aus Abbildung 5.1(b)	69
5.2	Ein aus Tabelle 5.1 abgeleitetes TDMA-Schema	71
7.1	Vor- und Nachteile des Nearest-Neighbours-Ansatzes	96
7.2	Vom Algorithmus verwendete Konstanten	104
8.1	Simulationsparameter	112

Abkürzungsverzeichnis

AODV	Ad-hoc On-demand Distance Vector routing, Seite 19
ARP	Address Resolution Protocol, Seite 110
CBTC	Cone-based Topology Control, Seite 57
CI	Contention Index, Seite 49
CSMA	Carrier Sense Multiple Access, Seite 11
CSMA/CA	Carrier Sense Multiple Access with Collision Avoidance, Seite 11
CTS	Clear To Send, Seite 13
DCF	Distributed Coordination Function, Seite 11
DIFS	DCF Interframe Space, Seite 12
DSR	Dynamic Source Routing, Seite 109
EIFS	Extended Interframe Space, Seite 12
ETSI	European Telecommunications Standards Institute, Seite 1
IFS	Interframe Space, Seite 12
ITU	International Telecommunication Union, Seite 1
LINT	Local Information, No Topology, Seite 49
MAC	Medium Access Control, Seite 110
MMLDP	Mobile Mesh Link Discovery Protocol, Seite 15
PCF	Point Coordination Function, Seite 11
PRN	Packet Radio Network, Seite 1
RERR	Router Error, Seite 20
RNG	Relative Neighbourhood Graph, Seite 53
RREP	Route Reply, Seite 19
RREQ	Route Request, Seite 19
RTS	Request To Send, Seite 13
SIFS	Short Interframe Space, Seite 12
SINR	Signal-to-Interference-and-Noise Ratio, Seite 6
SLF	Saturation Largest First, Seite 70
TDMA	Time Division Multiple Access, Seite 23
WLAN	Wireless Local Area Network, Seite 1

1 Einleitung

Das Gebiet der Mobilkommunikation ist derzeit noch in zwei relativ klar getrennte Teilbereiche untergliedert: In der „Telekommunikationswelt“ ermöglichen Mobilfunknetze die drahtlose Telefonie in weitläufigen Gebieten. Dieser Bereich wird beherrscht durch Anbieter von Telekommunikationsdienstleistungen, durch Hersteller von Komponenten für Mobilfunknetze und Mobiltelefonen und durch Organisationen wie die ETSI (European Telecommunications Standards Institute) und die ITU (International Telecommunication Union). In der „Rechnernetzwelt“ dagegen bieten Technologien zur drahtlosen Datenkommunikation die Möglichkeit, Personal-Computer (PCs) über kürzere Distanzen drahtlos zu lokalen Netzen (engl. *Wireless Local Area Networks*, WLANs) zusammenzuschließen. Dieser Bereich ist hauptsächlich geprägt durch die Aktivitäten der Hersteller von PC-Hardware und von Standardisierungsgremien wie der IEEE (Institute of Electrical and Electronics Engineers).

Gegenwärtig ist bereits zu beobachten, dass diese beiden Bereiche beginnen, zusammenzuwachsen. Obwohl Mobilfunknetze inzwischen auch in zunehmendem Maße zur Datenkommunikation verwendet werden, stellen Mobilfunkanbieter neue WLAN-Infrastruktur auf, um in räumlich begrenzten Gebieten (so genannten „Hotspots“) Datendienste mit wesentlich höherer Bandbreite anzubieten, als es mit Mobilfunknetzen auf absehbare Zeit möglich sein wird. Auch die Endgeräte aus beiden Welten werden immer ähnlicher: Moderne Mobiltelefone sind inzwischen so leistungsstark, wie es vor etwas über einer Dekade PCs üblicherweise waren, und es ist völlig gängig, Anwendungen wie Termin- und Adressverwaltung auf dem Mobiltelefon zu nutzen. Auch werden immer mehr Mobiltelefone mit WLAN-Technologien ausgestattet, die zurzeit ausschließlich zur unmittelbaren Kommunikation zwischen zwei Endgeräten verwendet werden (z.B. dem Austausch digitaler Visitenkarten). Zum anderen werden PCs in immer kleineren Größen hergestellt: Personal-Digital-Assistants (PDAs) sind bereits genau so klein, wie es große Mobiltelefone sind, die neben der Telefonie auch für andere Anwendungen gedacht sind.

Es besteht guter Grund zu der Annahme, dass sich diese Verschmelzung der beiden Bereiche in Zukunft fortsetzen wird. So wie heute ein Großteil der Bevölkerung der Industrieländer im Besitz von Mobiltelefonen ist, werden in Zukunft in hoher Dichte Geräte vorhanden sein, die eine Kombination aus Mobiltelefon und PC mit WLAN-Technologie sind. Vor diesem Hintergrund wird ein älteres Forschungsgebiet aus dem Bereich der Rechnernetze in neuem Licht interessant: Die Packet-Radio-Networks (PRNs) [Kah77].

PRNs bestehen aus Rechnern, die ohne Zuhilfenahme zusätzlicher Infrastruktur drahtlos miteinander verbunden sind. Dazu wird eine Broadcast-Funktechnologie verwendet; die Stationen nutzen die zur Kommunikation verwendeten Ressourcen (Frequenzbereiche) also ge-

meinsam. (Damit sind Richtfunkstrecken ausgeschlossen, die beispielsweise Mobiltelefonie-Basisstationen mit einem Kernnetz verbinden, da Richtfunk für mobile Geräte nicht anwendbar ist und in seinen Charakteristiken eher einer drahtgebundenen Übertragungstechnologie als einer Broadcast-Funktechnologie gleicht.) Können zwei Stationen nicht direkt miteinander kommunizieren, so fungieren in einem PRN andere Stationen als Relais. Während eine Broadcast-Funktechnologie in fast allen derzeit existierenden Netzen ausschließlich dazu dient, ein Endgerät an eine Infrastruktur anzubinden, sind PRNs echte drahtlose Multihop-Netze, d.h. Netze, in denen Datenpakete über potenziell beliebig viele Relais drahtlos weitergeleitet werden.

Ist eine genügend hohe Dichte von Geräten mit WLAN-Technologie vorhanden, ist auch hier der Gedanke naheliegend, mit Hilfe von Multihop-Kommunikation den Datenaustausch zwischen Geräten zu ermöglichen, die keinen direkten Funkkontakt zueinander haben. Während PRNs aber weitgehend statische Netze sind, deren einzige Dynamik dadurch zustande kommt, dass in vergleichsweise großen zeitlichen Abständen Geräte zum Netz hinzugefügt oder aus dem Netz entfernt werden, werden die drahtlosen Multihop-Netze der Zukunft durch die Mobilität der Geräte wesentlich dynamischer sein. Diese Netze werden daher meist als *mobile Ad-hoc-Netze* bezeichnet, um auszudrücken, dass die beteiligten Geräte mobil sein können und dass eine manuelle Konfiguration wegen der hohen Dynamik unpraktikabel ist.

In den letzten Jahren hat sich ein Großteil der Arbeiten zu mobilen Ad-hoc-Netzen mit der naheliegendsten Herausforderung auseinander gesetzt, dem selbstkonfigurierenden Routing. Ein anderer Themenbereich, dessen Ursprung schon in der Forschung zu PRNs zu finden ist, ist die Sendeleistungsregelung der beteiligten Geräte. Während eine Vielzahl aktueller Arbeiten eine dynamische Regelung auf möglichst geringe Sendeleistungen als Mittel zur verlängerten Haltbarkeit der Batterien tragbarer, kleiner Geräte verwendet, ist die ursprüngliche Motivation auch für mobile Ad-hoc-Netze von großem Interesse, nämlich der Einfluss der Sendeleistungsregelung auf die Kapazität des Netzes.

Intuitiv kann man sich vorstellen, dass die Sendeleistungen der Geräte nicht zu gering sein dürfen, weil das Netz sonst in mehrere Teile zerfällt, zwischen denen keine Kommunikationspfade etabliert werden können. Zu hohe Sendeleistungen haben hingegen den Nachteil, dass sich eine große Zahl von Geräten das gemeinsame Kommunikationsmedium teilen muss, das damit also nur während eines geringen Anteils der Gesamtzeit einem einzelnen Nutzer zur Verfügung steht. Aus diesem Grund kann es zur Steigerung der Netzkapazität sinnvoll sein, die Geräte mit geringerer Sendeleistung zu betreiben und damit eine höhere Zahl zeitlich überlappend stattfindender Übertragungen zuzulassen, d.h. den *Spatial-Reuse* zu erhöhen.

Deshalb sind Verfahren gefragt, die den am Netz beteiligten Geräten Sendeleistungen so zuweisen, dass den Nutzern eine möglichst hohe Kapazität zur Verfügung steht. In Ermangelung einer zentralen Instanz, die Daten zur Netztopologie sammeln und auswerten könnte, ist dies kein einfaches Unterfangen, erst recht nicht, wenn keine besonderen Anforderungen an die verwendete Hardware und das geplante Einsatzszenario gestellt werden sollen.

Die vorliegende Arbeit leistet mehrere Beiträge zu diesem Themengebiet: Zunächst wird der Zusammenhang zwischen der Topologie und der Kapazität eines Netzes genau beleuchtet, um

entscheiden zu können, welche Topologieklassen eine hohe Netzkapazität ermöglichen. Dann wird ein Verfahren zur dezentralen Sendeleistungsregelung vorgestellt, das die Netztopologie einem bestimmten Graphtypen annähert, von dem zuvor gezeigt wurde, dass er ein hohes Kapazitätspotenzial aufweist und der sich auch in praktischer Hinsicht gut zur dezentralen Sendeleistungsregelung eignet. Im Gegensatz zu bereits existierenden Verfahren zur Anpassung der Netztopologie an diesen Graphtypen ist das Verfahren praktikabel in dem Sinn, dass es keine besonderen Eigenschaften der drahtlosen Hardware erfordert (außer der Möglichkeit zur Sendeleistungsregelung an sich) und dass es auch in Netzen anwendbar ist, die wie oben beschrieben eine gewisse Dynamik aufweisen.

Die Arbeit ist folgendermaßen strukturiert: In Kapitel 2 werden die zum Verständnis der folgenden Kapitel benötigten Grundlagen eingeführt. In Kapitel 3 erfolgt eine genaue Darlegung der Annahmen, die der vorliegenden Arbeit zugrunde liegen, und der verfolgten Ziele. Kapitel 4 gibt einen Überblick über verwandte Arbeiten und erläutert, wie sich die vorliegende Arbeit dagegen abgrenzt. Kapitel 5 stellt eigene Verfahren vor, die zur Kapazitätsabschätzung der Topologien drahtloser Netze dienen. Anhand dieser Verfahren wird in Kapitel 6 die Kapazität unterschiedlicher in Kapitel 4 vorgestellter Topologieklassen bewertet. In Kapitel 7 wird ein neues verteiltes Verfahren zur Sendeleistungsregelung in drahtlosen Netzen präsentiert. In Kapitel 8 wird die Eignung des Verfahrens und seiner unterschiedlichen Varianten in Simulationen unter verschiedenen Bedingungen überprüft. In Kapitel 9 werden schließlich die Kernpunkte der Arbeit zusammengefasst und ein Ausblick auf offene Fragestellungen gegeben.

2 Grundlagen

2.1 Drahtlose Datenübertragung

Drahtlose Datenübertragung wird heutzutage praktisch ausschließlich mittels elektromagnetischer Wellen durchgeführt. Meist werden dazu Radio- oder Mikrowellen eingesetzt, teilweise aber auch Infrarotlicht. Letzteres hat jedoch den Nachteil, dass direkter Sichtkontakt und eine relativ geringe räumliche Distanz zwischen Sender und Empfänger erforderlich sind. Dies wird von Benutzern jedoch als unkomfortabel empfunden, weshalb Mobiltelefone gegenwärtig immer häufiger zusätzlich mit der Bluetooth-Technologie zur Funkkommunikation über kurze Distanzen ausgestattet werden. Weil die in dieser Arbeit betrachteten Szenarien unter solchen Voraussetzungen auch kaum zu realisieren sind, wird hier ausschließlich drahtlose Datenübertragung mit Hilfe von Radiowellen untersucht.

2.1.1 Antennen

Die verwendeten Antennen haben einen starken Einfluss auf die Übertragungseigenschaften von Funksignalen. Verschiedene Antennentypen strahlen das Signal unterschiedlich stark gerichtet ab. Omnidirektionale Antennen strahlen das Signal in alle Richtungen mit annähernd gleich starker Leistung ab, während Richtfunkantennen das Signal strahlförmig aussenden. Im ersten Fall ist die Dämpfung des Signals mit zunehmender Distanz relativ hoch, jedoch weitgehend unabhängig von der Richtung des Empfängers. Im zweiten Fall nimmt die Empfangsleistung mit steigender Distanz des Empfängers vom Sender kaum ab, der Empfänger muss jedoch genau im Strahl des Senders liegen, sonst ist ein Empfang gar nicht möglich. Zwischen diesen beiden Extremen gibt es Zwischenstufen, beispielsweise Antennen, die einen Sektor mit einem bestimmten Winkel bedienen.

Je stärker das Signal gebündelt wird, desto günstiger sind zwar die Übertragungseigenschaften, da das Signal den Empfänger mit stärkerer Leistung erreicht, während es an anderen Geräten schwächer ist und damit weniger Interferenz verursacht. Umso komplizierter ist es jedoch, dafür zu sorgen, dass der Empfänger in dem Bereich der Bündelung liegt, vor allem wenn die Geräte ständig in Bewegung sein können. Daher ist davon auszugehen, dass in den betrachteten Szenarien in erster Linie omnidirektionale Antennen eingesetzt werden.

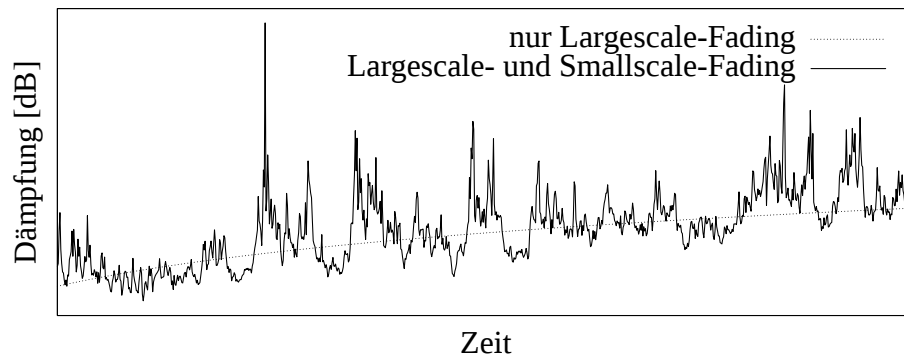


Abbildung 2.1: Veranschaulichung von durch Largescale- und Smallscale-Fading verursachter Dämpfung

2.1.2 Bedingungen für eine erfolgreiche Datenübertragung

Es existieren diverse Möglichkeiten, um Daten in einem Funksignal zu kodieren, und es würde über den Rahmen der vorliegenden Arbeit hinausgehen, technische Details zu beschreiben. Relevant für das Verständnis der folgenden Kapitel ist lediglich, dass ein Funksignal nur dann fehlerfrei dekodiert werden kann, wenn es den Empfänger mit einer gewissen absoluten Mindeststärke erreicht, dem sog. *Receive-Threshold*. Dieser Schwellwert hängt mit der Robustheit der Kodierung und damit mit der Datenrate zusammen: Es entspricht der Intuition, dass ein Signal mit robusterer Kodierung auch bei geringerer Empfangsstärke noch fehlerfrei dekodiert werden kann, dass es dafür aber weniger Bits pro Zeiteinheit enthalten kann.

Für einen fehlerfreien Empfang ist es außerdem erforderlich, dass keine zu starken Störungen im gleichen Frequenzbereich vorliegen, dass der *Signal-to-Interference-and-Noise-Ratio* (SINR) also einen gewissen Schwellwert nicht unterschreitet. Auch hier ist der Zusammenhang zur Robustheit der Kodierung dadurch gegeben, dass Signale mit robusterer Kodierung auch bei geringerem SINR noch fehlerfrei dekodiert werden können.

Moderne WLAN-Hardware beherrscht mehrere Kodierungen, so dass situationsabhängig eine angemessene Wahl erfolgen kann. Für Links mit vergleichsweise geringer Dämpfung bietet sich eine weniger robuste Kodierung mit hoher Datenrate an. Für Links mit stärkerer Dämpfung muss hingegen eine robustere Kodierung verwendet und somit eine geringere Datenrate in Kauf genommen werden.

2.1.3 Die Ausbreitung von Funkwellen

Der folgende Abschnitt, der die Modellierung der Ausbreitungseigenschaften von Funksignalen erläutert, lehnt sich stark an [Rap96] und [TSRK04] an. Sämtliche angegebenen Formeln sind diesen Büchern entnommen.

Die Signalstärke an einem Empfänger errechnet sich durch die Sendeleistung abzüglich der *Dämpfung* bzw. des *Pathloss*. Effekte, die diese Dämpfung hervorrufen, werden unter dem Begriff des *Fading* zusammengefasst. Das *Largescale-Fading* verursacht den Anteil der Dämpfung, der unter anderem von der Distanz zwischen Sender und Empfänger abhängt, der aber unabhängig von der Zeit und von kleinen räumlichen Verschiebungen ist, während das *Smallscale-Fading* kurzfristige Schwankungen der Signalstärke verursacht. Abbildung 2.1 veranschaulicht diese beiden Komponenten.

Im folgenden Abschnitt 2.1.3.1 wird die Modellierung des Largescale-Fading erläutert, und im darauf folgenden Abschnitt 2.1.3.2 wird auf die Modellierung des Smallscale-Fadings eingegangen.

2.1.3.1 Largescale-Fading

Es gibt verschiedene Modelle zur Berechnung des Largescale-Pathloss, die auf speziellen Annahmen über die Umgebung basieren. Das Friis'sche *Free-Space-Modell* etwa modelliert die Ausbreitung von Radiowellen im vollständig freien Raum, so dass diese den Empfänger ausschließlich entlang der direkten Sichtlinie erreichen. Somit ist die Dämpfung in Abhängigkeit der Sender-Empfänger-Distanz d proportional zu d^2 (also zur Oberfläche einer Kugel mit Radius d). Das *Two-Ray-Ground-Reflection-Modell* berücksichtigt zwei Ausbreitungspfade, nämlich die direkte Sichtlinie und einen Pfad, der durch Reflektion am Boden entsteht. Die Dämpfung ist durch die möglicherweise phasenverschobene reflektierte Komponente potenziell stärker; der Largescale-Pathloss ist hier proportional zu d^4 .

Ein generisches Modell zur Berechnung des Largescale-Pathloss ist das *Log-Distance-Modell*, das auf der Erkenntnis beruht, dass die mittlere Signalstärke in dB mit steigender Distanz logarithmisch fällt. Die mittlere Dämpfung über eine Distanz d ist damit (in dB) gegeben durch die mittlere Dämpfung über eine Referenzdistanz d_0 (in dB) und den Pathloss-Exponenten n :

$$\overline{PL}(d) = \overline{PL}(d_0) + 10 \cdot n \cdot \log(d/d_0)$$

Damit ist die mittlere Dämpfung proportional zu d^n , speziellere Modelle wie die eingangs erwähnten sind hier also durch eine geeignete Wahl der Parameter enthalten.

Da insbesondere in urbaner Umgebung die Dämpfung nicht nur von der Distanz an sich abhängt, sondern je nach konkreter Lage von Sender und Empfänger stark schwanken kann, wird der Largescale-Pathloss im *Lognormal-Shadowing-Modell* zusätzlich durch eine normalverteilte Variable X_σ mit Mittelwert 0 und Standardabweichung σ bestimmt:

$$PL(d) = \overline{PL}(d) + X_\sigma$$

Da dieses Modell in aktuellen Forschungsarbeiten häufig zur Modellierung des Smallscale-Fadings genutzt wird, sei an dieser Stelle betont, dass es sich hierbei um ein Largescale-Fading-Modell handelt, das also die Verteilung der mittleren Signalstärke in Abhängigkeit von der Sender-Empfänger-Distanz beschreibt. Befinden sich Sender und Empfänger aber an

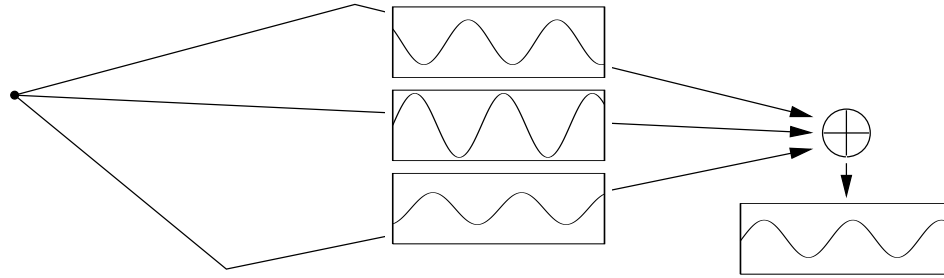


Abbildung 2.2: Veranschaulichung der Konsequenzen von Multipfad-Ausbreitung

festgelegten Positionen, ändert sich die mittlere Signalstärke nicht. Eigene Experimente haben gezeigt, dass die Verwendung als Smallscale-Fading-Modell deshalb ungünstig ist, weil unrealistischerweise auch Signalstärken, die deutlich über dem Mittelwert liegen, vergleichsweise häufig auftreten, und diese sich als weit reichende Störungen bemerkbar machen.

2.1.3.2 Smallscale-Fading

Das Smallscale-Fading wird in erster Linie durch Verstärkungs- und Auslöschungseffekte verursacht, die dadurch entstehen, dass sich am Empfänger unterschiedlich verzögerte Versionen des ausgesendeten Signals überlagern (siehe Abbildung 2.2). Die einzelnen Komponenten, aus denen sich das Gesamtsignal am Empfänger zusammensetzt, werden auch *Multipfad-Komponenten* genannt, weil sie den Empfänger über verschiedene Ausbreitungswege erreichen. Während eine Komponente den Empfänger gegebenenfalls entlang der direkten Sichtlinie (engl. *Line-Of-Sight-Path*) erreichen kann (abhängig von der Umgebung und der Lage von Sender und Empfänger darin), kommt es in allen praktisch relevanten Szenarien vor, dass Kopien des Signals durch Reflektion an Objekten in der Umgebung über zusätzliche Pfade mit unterschiedlichen Verzögerungen den Empfänger erreichen. Da sich diese Multipfad-Komponenten in ihrer Phase unterscheiden, können sie sich gegenseitig verstärken oder abschwächen.

Werden Phase ϕ_k und die Amplitude a_k einer Multipfad-Komponente k als Polarkoordinaten eines komplexen Werts

$$a_k \cdot e^{i\phi_k} = x_k + iy_k$$

aufgefasst, sind die Eigenschaften der Überlagerung mehrerer Multipfad-Komponenten am Empfänger durch die Polarkoordinaten der Summe

$$r = \sum_k a_k \cdot e^{i\phi_k} = \sum_k x_k + i \sum_k y_k$$

beschrieben.

Unter der Annahme, dass hinreichend viele Multipfad-Komponenten vorhanden sind, deren Amplituden identisch verteilt sind und deren Phasen gleichverteilt sind, sind $\sum_k x_k$ und $\sum_k y_k$ nach dem zentralen Grenzwertsatz normalverteilt um den Erwartungswert 0. Die Amplitude

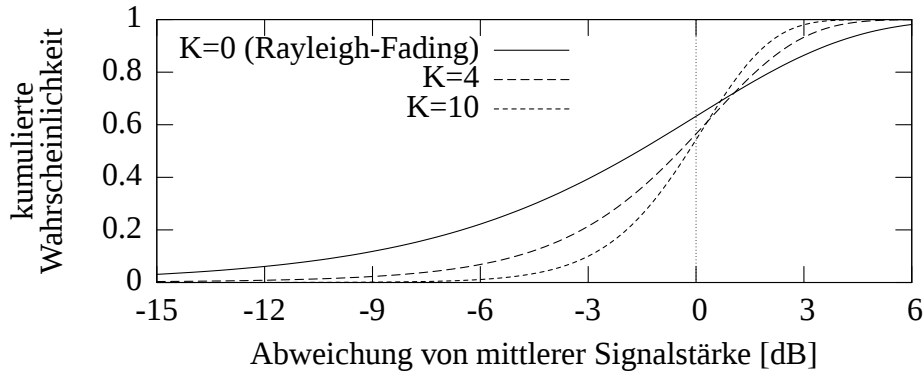


Abbildung 2.3: Signalstärkeschwankungsverteilungen für verschiedene Rice-Faktoren

des zusammengesetzten Signals $|r|$ folgt dann einer *Rayleigh-Verteilung*, deren Dichte gegeben ist durch

$$P(x) = \begin{cases} \frac{x}{\sigma^2} e^{-\frac{x^2}{2\sigma^2}} & \text{für } x \geq 0, \\ 0 & \text{für } x < 0. \end{cases}$$

Dies ist insbesondere dann jedoch nicht der Fall, wenn eine dominante Komponente vorhanden ist, beispielsweise diejenige, die den Empfänger direkt entlang der Sichtlinie erreicht. In diesem Fall ist der Erwartungswert wenigstens einer der beiden Normalverteilungen größer als 0 und die Amplitude des zusammengesetzten Signals folgt einer *Rice-Verteilung*, deren Dichte gegeben ist durch

$$P(x) = \begin{cases} \frac{x}{\sigma^2} e^{-\frac{x^2+A^2}{2\sigma^2}} I_0\left(\frac{Ax}{\sigma^2}\right) & \text{für } x \geq 0, \\ 0 & \text{für } x < 0. \end{cases}$$

Der Parameter $A \geq 0$ gibt die Amplitude der dominanten Komponente an, so dass die Rice-Verteilung mit $A = 0$ identisch zur Rayleigh-Verteilung ist, und

$$I_0(x) = \frac{1}{\pi} \int_0^\pi e^{x \cos \theta} d\theta$$

bezeichnet die modifizierte Bessel-Funktion erster Art der Ordnung 0.

Zwischen der Stärke eines Signals und dem Quadrat der Amplitude besteht ein linearer Zusammenhang. Wird die Amplitude durch eine Zufallsvariable X beschrieben, besteht also ein linearer Zusammenhang zwischen der mittleren Signalstärke, die den Gesetzen des Largescale-Fadings unterliegt, und dem Erwartungswert des Quadrats der Amplitude $E(X^2)$. Somit beschreibt die Zufallsvariable $X^2/E(X^2)$ die Schwankungen um die mittlere Signalstärke. Ist X Rice-verteilt, so ist $E(X^2) = A^2 + 2\sigma^2$, und die Zufallsvariable $X^2/E(X^2)$ ist durch den in der Literatur häufig als *Rice-Faktor* bezeichneten Parameter $K = A^2/(2\sigma^2)$ bereits vollständig festgelegt, der das Verhältnis der Signalstärken der dominanten Komponente

und Überlagerung der restlichen Multipfad-Komponenten angibt. Rayleigh-Fading ist wiederum als Spezialfall $K = 0$ enthalten. Abbildung 2.3 zeigt die Verteilungen der Signalstärke-Schwankungen für unterschiedliche K . Wie dort zu erkennen ist, sind die Schwankungen um die mittlere Signalstärke in der logarithmischen dB-Einheit nicht symmetrisch wie im Lognormal-Shadowing-Modell (siehe Abschnitt 2.1.3.1), vielmehr treten wesentlich häufiger tiefe Einbrüche der Signalstärke auf, die auch im Beispiel zu sehen sind, das in Abbildung 2.1 dargestellt ist.

Die obigen Ausführungen beziehen sich auf den statischen Fall, in dem Sender und Empfänger sich nicht bewegen. Wie bereits erläutert, ändern sich die Eigenschaften einer Multipfad-Komponente durch geringfügige Bewegung des Senders, des Empfängers, oder anderer Objekte in der Umgebung. Insbesondere ist die Phase einer Multipfad-Komponente abhängig von der Distanz, die das Signal entlang des Pfads zurücklegt. Durch Veränderung der Pfadlänge um eine Distanz von maximal einer Wellenlänge können bereits sämtliche Phasenverschiebungen zwischen 0 und 2π durchlaufen werden. Ändert sich der Längenunterschied zwischen zwei Multipfad-Komponenten um eine halbe Wellenlänge, wird aus einem Verstärkungseffekt ein Auslöschungseffekt und umgekehrt. Angesichts der Tatsache, dass in der drahtlosen Kommunikation eher kleine Wellenlängen verwendet werden (die Wellenlänge eines 2,4GHz-Signals beträgt beispielsweise 12,5cm) ist klar, dass die Empfangsstärke eines Signals sich schon durch geringfügige Bewegungen stark verändern kann.

Die beschriebene Modellierung des Smallscale-Fadings setzt voraus, dass alle relevanten Frequenzanteile des Signals in gleicher Weise beeinflusst werden, dass also ein sog. *flaches* Fading vorliegt. Dies ist jedoch nicht zwingend der Fall. Je größer die Differenz zwischen zwei Frequenzen ist, desto schwächer korrelieren deren Amplituden, so dass breitbandigere Signale mit höherer Wahrscheinlichkeit einem *frequenzselektivem* Fading unterliegen. Die sog. *Kohärenz-Bandbreite*, die angibt, wie groß der Unterschied zwischen zwei Frequenzen sein darf, damit die Amplituden beider Frequenzen am Empfänger hinreichend stark korrelieren, ist u.a. abhängig von der Streuung der Verzögerungen entlang verschiedener Multipfad-Komponenten. Ein genauer Zusammenhang existiert auf dieser abstrakten Ebene jedoch nicht, sondern hängt vom Zusammenspiel verschiedener Faktoren ab, so dass hier häufig nur Simulation mit detaillierter Modellierung der physikalischen Ebene Aufschluss geben kann.

2.2 Die Verbindungsschicht in drahtlosen Netzen

Die *Verbindungsschicht* (oder der *Data-Link-Layer*) ist die Ebene 2 des OSI-Referenzmodells [Int94]. Sie regelt die Kommunikation zwischen benachbarten Geräten und nutzt dafür die Bit-übertragungsschicht, die die physikalische Übertragung einzelner Bits ermöglicht. Die zwei Hauptaufgaben der Verbindungsschicht sind die Regelung des Medienzugangs und ggf. die Fehlererkennung oder -korrektur.

Es gibt zwei grundlegend verschiedene Arten, den Zugang zu einem gemeinsam genutzten Medium zur drahtlosen Datenübertragung zu regeln, bei dem es sich im hier betrachteten Zu-

sammenhang um bestimmte Radiofrequenzen in einem gewissen räumlichen Bereich handelt.

Single-Channel-Technologien verwenden den gesamten zur Verfügung stehenden Frequenzbereich als Übertragungskanal. Deshalb können Übertragungen nur dann zeitlich überlappend stattfinden, wenn die räumliche Trennung zwischen ihnen eine so starke Dämpfung der Funksignale bewirkt, dass das SINR bei den Empfängern genügend hoch ist. Diese Art der mehrfachen, zeitlich überlappend stattfindenden Nutzung des Kanals wird auch als *Spatial-Reuse* bezeichnet.

Multi-Channel-Technologien teilen die Ressourcen in mehrere Kanäle auf, die zeitgleich für verschiedene Übertragungen genutzt werden können, ohne dass diese sich gegenseitig stören. Die Kenntnis technischer Möglichkeiten zur Realisierung solcher Technologien sind für das Verständnis der vorliegenden Arbeit nicht notwendig, so dass dieser Themenbereich hier nicht behandelt wird.

Der Vorteil von Multi-Channel-Technologien liegt auf der Hand: Da unterschiedliche Sender-Empfänger-Paare verschiedene Kanäle verwenden, können zu einem Zeitpunkt mehrere Übertragungen stattfinden, ohne sich gegenseitig zu stören. Der Nachteil ergibt sich durch die technisch bedingte Einschränkung, dass jedes Gerät zu einem bestimmten Zeitpunkt nur in einem Kanal aktiv sein kann, so dass die Notwendigkeit besteht, die Geräte zu synchronisieren. In Multi-Channel-Netzen ist dies ein Problem, das theoretisch zwar von polynomieller Komplexität ist ([HS88]), praktisch jedoch außerordentlich schwierig ist, wie beispielsweise die Vielzahl von Arbeiten zum Scheduling in Bluetooth-Netzen zeigt (z.B. [BFK⁺02]).

Mit einer Single-Channel-Technologie lassen sich drahtlose Netze dagegen wesentlich unkomplizierter realisieren. Die vorliegende Arbeit konzentriert sich auf solche Netze mit gemeinsam genutztem Übertragungskanal (wie schon im Titel und auch in Abschnitt 3.1 nochmals festgehalten ist), so dass hier die einzige Single-Channel-Technologie vorgestellt wird, die im kommerziellen Bereich derzeit von praktischer Relevanz ist, nämlich IEEE 802.11. Die heute ebenfalls stark verbreitete Multi-Channel-Technologie Bluetooth [Blu04], die in einer vorangegangenen Version als IEEE 802.15.1 standardisiert wurde [IEE02], wird nicht näher erläutert.

2.2.1 IEEE 802.11

Derzeit ist der Begriff WLAN im Alltagsgebrauch synonym mit den Standards der IEEE-802.11-Familie, weil praktisch keine andere Technologie zum Aufbau drahtloser lokaler Computernetze eingesetzt wird. Der Basis-Standard [IEE99] spezifiziert mehrere alternative physikalische Schichten (sowohl für Funk- als auch für Infrarot-Datenübertragung) und eine darauf aufbauende Medienzugangsschicht mit mehreren Betriebsmodi.

IEEE 802.11 ist eine Single-Channel-Technologie. Sie ist deshalb so konzipiert, dass andere Sender eine laufende Übertragung möglichst selten stören. Dies geschieht durch die Kontrolle des Medienzugangs mittels Carrier-Sense-Multiple-Access mit Collision-Avoidance (CSMA/CA).

Dazu werden zwei alternative Verfahren spezifiziert, und zwar die *Distributed-Coordination-Function* (DCF) und die *Point-Coordination-Function* (PCF). Unter Verwendung der DCF sind alle Geräte gleichberechtigt, während gemäß der PCF eine zentrale Instanz (der *Point-Coordinator*) festlegt, in welchen Zeiträumen welches Gerät senden darf bzw. wann Random-Access-Phasen sind, in denen die Geräte, wie auch bei der DCF, um den Zugang zum gemeinsamen Übertragungsmedium konkurrieren. Für die DCF unterscheidet der Standard weiterhin zwischen einem *Infrastruktur-Modus*, der dem Access Point eine zentrale Rolle zuweist (dieser ist eine in der Regel fest installierte Bridge, die in vielen Anwendungsszenarien zur Anbindung an andere Netze benötigt wird), und dem *Ad-hoc-Modus*, der auf zentrale Elemente verzichtet. Alle weiteren Ausführungen beziehen sich nur auf den Ad-hoc-Modus der DCF, weil die vorliegende Arbeit sich mit drahtlosen Netzen befasst, die auch ohne Infrastruktur funktionsfähig sind.

Um Rahmen unterschiedlich zu priorisieren, spezifiziert IEEE 802.11 mehrere *Interframe-Spaces* (IFS), die festlegen, welche Zeitspanne zwischen dem Ende der vorangegangenen und dem Beginn der nächsten Übertragung liegen muss. Je kürzer der IFS, desto höher ist also die Priorität einer Übertragung.

Der kürzeste IFS, der *Short-IFS* (SIFS), sorgt dafür, dass Rahmensequenzen möglichst nicht unterbrochen werden, die erfordern, dass mehrere Übertragungen unmittelbar aufeinander folgen. Zu diesen Mechanismen zählen der Quittierungsmechanismus für Unicast-Übertragungen und der RTS/CTS-Austausch, die weiter unten erläutert werden.

Der nächstlängere *DCF-IFS* (DIFS) gilt für den Beginn einer neuen Rahmensequenz (die auch aus einem einzelnen Rahmen bestehen kann). Der CSMA/CA-Mechanismus beinhaltet hier außerdem, dass zusätzlich ein randomisierter Backoff stattfindet, wenn das Medium zu dem Zeitpunkt belegt ist, an dem der erste Rahmen gesendet werden soll. Dadurch wird die Wahrscheinlichkeit verringert, dass mehrere Geräte nach dem Ende einer Übertragung die gleiche Wartezeit einhalten, dann zeitgleich beginnen zu senden und so eine Kollision verursachen.

Der *Extended-IFS* (EIFS) ist der längste IFS und wird von Geräte eingehalten, die einen fehlerhaften Rahmen empfangen. Dies soll einem anderen Gerät ermöglichen, den möglicherweise korrekt empfangenen, an ihn adressierten Rahmen wie unten beschrieben zu quittieren. Die Länge des EIFS ist daher genau die zur Quittierung eines Rahmens benötigte Zeit zuzüglich eines DIFS.

Wegen der hohen Fehleranfälligkeit drahtloser Datenübertragungen schreibt IEEE 802.11 einen Quittierungsmechanismus für Unicast-Datenübertragungen vor: Ein korrekt empfangenes Datenpaket wird vom Empfänger nach einem SIFS quittiert. Empfängt der Sender innerhalb des *ACK-Timeout* (einem SIFS zuzüglich der Dauer eines Quittungsrahmens) keine Quittung, wird die Übertragung deshalb nach einer zufälligen Wartezeit wiederholt. Nach einer bestimmten Anzahl erfolgloser Übertragungsversuche beendet der Sender den Vorgang mit einer Fehlermeldung an die höheren Protokollschichten.

Der beschriebene CSMA/CA-Mechanismus ist gut geeignet, um Kollisionen zu vermeiden, wenn die Geräte immer in der Lage sind, zu erkennen, dass der Übertragungskanal belegt ist.

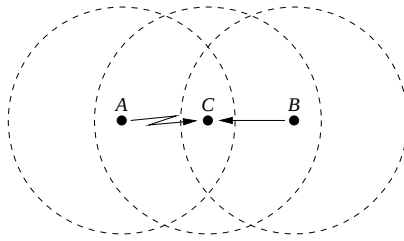


Abbildung 2.4: Das Hidden-Node-Problem

Sind die Geräte jedoch über einen größeren räumlichen Bereich verteilt, so kann die Dämpfung zwischen zwei Geräten A und B so stark sein, dass A ein Signal, das B sendet, nicht wahrnehmen kann. Da das Übertragungsmedium aus der Sicht von A frei ist, kann auch A unmittelbar mit einer Übertragung beginnen, obwohl das Signal von A einen Empfänger C der Übertragung von B noch mit hinreichender Stärke erreichen kann, um den Empfang zu stören. Dieses *Hidden-Node-Problem* ist in Abbildung 2.4 skizziert.

Um diese Situation zu verhindern, spezifiziert IEEE 802.11 zusätzlich zum physikalischen Carrier-Sensing ein *virtuelles Carrier-Sensing*, das im beschriebenen Beispiel das Gerät A darüber informieren soll, dass das Medium belegt ist, obwohl A physikalisch keine laufende Übertragung wahrnehmen kann. Dazu dient der *RTS/CTS-Mechanismus* (Request To Send / Clear To Send). Vor der Übertragung eines längeren Datenpakets verschickt ein Sender ein RTS, das das beabsichtigte Zielgerät und die Dauer der geplanten Übertragung enthält. Ist das Zielgerät empfangsbereit, beantwortet es das RTS mit einem CTS, das ebenfalls die im RTS enthaltene Übertragungsdauer der folgenden Datenübertragung enthält. Alle Geräte, die ein CTS empfangen, erachten das Medium für den Zeitraum der folgenden Übertragung als belegt.

Obwohl der RTS/CTS-Mechanismus das Hidden-Node-Problem abschwächen kann, ist er nicht in der Lage, es vollkommen zu lösen. Zum einen ist er nur bei Unicast-Übertragungen anwendbar, Broadcasts sind hingegen ungesichert. Außerdem kann auch ein CTS selbst durch das Hidden-Node-Problem nicht immer korrekt empfangen werden. Und letztendlich kann ein Signal eines Senders auch einen Empfänger stören, wenn die Dämpfung zwischen beiden Geräten schon zu stark ist, als dass ein fehlerfreier Empfang eines CTS-Pakets mit akzeptabler Wahrscheinlichkeit noch möglich wäre.

Um das Vorhandensein und die Konfiguration des drahtlosen Netzes anzuzeigen und die teilnehmenden Geräte zu synchronisieren, werden in festen Zeitintervallen *Beacon-Rahmen* generiert. Ein solcher Rahmen wird nach einem zufälligen Backoff versendet, der aus einem Fenster gewählt wird, das die doppelte Größe dessen hat, das aktuell für Datenrahmen verwendet wird. Die Geräte, die ein Beacon empfangen, bevor sie ihr eigenes versenden können, brechen den Backoff ab und werfen ihr Beacon. Damit versendet ein einzelnes Gerät im Regelfall also nur in unregelmäßigen Abständen Beacons; dieser Mechanismus ist damit nicht geeignet zur Linkerkennung wie in Abschnitt 2.2.2 beschrieben.

Die einzigen physikalischen Schichten, die in IEEE-802.11-Produkten tatsächlich verwendet

werden, sind die *Direct-Sequence-Spread-Spectrum*- und *Orthogonal-Frequency-Division-Multiplex-Schichten* (DSSS, OFDM). Sie ermöglichen durch die Verwendung unterschiedlicher Kodierungen verschiedene Datenraten (siehe Abschnitt 2.1), wobei nicht zwingend alle Geräte die gleiche Menge an Kodierungen unterstützen müssen. Deshalb muss für jede Unicast-Übertragung entschieden werden, welche Datenrate zu verwenden ist (für Broadcast-Übertragungen wird grundsätzlich die *Basic-Rate* verwendet, die alle Geräte beherrschen müssen, die an dem Netz teilnehmen).

IEEE 802.11 ist nicht in erster Linie für drahtlose Multihop-Netze konzipiert, und aktuelle Forschung zeigt, dass es einige Nachteile in einer solchen Umgebung gibt. So wird durch den RTS/CTS-Mechanismus zwar versucht, das Hidden-Node-Problem abzuschwächen bzw. zu beseitigen, daraus ergibt sich jedoch eine andere Schwierigkeit: Geräte, die ein CTS empfangen haben und das Medium deshalb als belegt erachten, antworten nicht auf RTS-Rahmen, die an sie adressiert sind. Dieses an sich sinnvolle Verhalten kann aber dazu führen, dass während der Übertragung eines längeren Datenrahmens so viele dieser vergleichsweise kurzen Anfragen unbeantwortet bleiben, dass der dafür vordefinierte Schwellwert überschritten wird und die Übertragung mit einer Fehlermeldung an die höheren Schichten abgebrochen wird. Routingprotokolle für drahtlose Multihop-Netze interpretieren eine solche Fehlermeldung so, als sei der betroffene Link weggefallen und versuchen in der Folge, neue Routen zu etablieren, was mit einem hohen Signalisierungsaufwand verbunden ist. Aus diesem Grund ist der Durchsatz, der mit TCP in drahtlosen Multihop-Szenarien erzielt werden kann, sehr gering ([XS01], [XS02], [FLZ⁺05]). Außerdem zeigen verschiedene Untersuchungen, dass IEEE 802.11 keine faire Aufteilung der Sendezeiten erlaubt, wenn die Geräte nicht alle in gegenseitiger Sendereichweite sind (siehe z.B. [LNG03], [LNG04], [SHS04]).

2.2.2 Drahtlose Links und Interferenz

Das im vorangegangenen Abschnitt vorgestellte Protokoll sieht (wie viele andere Protokolle auf Verbindungsschicht auch) an sich keine Mechanismen vor, um das Vorhandensein von Links festzustellen. Zunächst ist also kein Wissen darüber vorhanden, welche Geräte direkt miteinander kommunizieren können. Dieser Abschnitt beschreibt Verfahren, mit deren Hilfe dieses Wissen beschafft werden kann.

Wie in Abschnitt 2.4.1 genauer beschrieben wird, beruhen viele Arbeiten auf der Annahme, dass die Stärke eines empfangenen Signals ausschließlich von der Sendeleistung und der Entfernung zwischen Sender und Empfänger abhängt und mit steigender Distanz streng monoton fällt, so dass (in Abwesenheit von Interferenz) eine Datenübertragung zwischen zwei Geräten in einem statischen Netz entweder immer oder nie fehlerfrei stattfinden kann. Dies ist jedoch eine starke Vereinfachung, denn das Smallscale-Fading führt zu deutlichen Schwankungen der Signalstärke, selbst wenn sich Sender und Empfänger gar nicht oder nur geringfügig bewegen. Deshalb ist es wesentlich näher an der Realität, einen Link durch eine Verlustwahrscheinlichkeit zu charakterisieren.

Die einfache Modellierung *binärer Linkzustände* ist geeignet, um abstrakte Untersuchungen

durchzuführen. Beim Design von Protokollen für drahtlose Netze sollte aber berücksichtigt werden, dass Links mit beliebigen Verlustwahrscheinlichkeiten zwischen 0 und 1 behaftet sein können. Gerade in Arbeiten zu drahtlosen Multihop-Netzen wird implizit jedoch sehr häufig davon ausgegangen, dass Links grundsätzlich zuverlässig sind oder dass ein Mechanismus existiert, der unzuverlässige Links ausblendet. Dies trifft beispielsweise auch für das in Abschnitt 2.3.4 beschriebene AODV-Protokoll zu. Beim Empfang einer RREQ-Nachricht, die als Broadcast-Rahmen an alle Geräte in Sendereichweite verschickt wird, wird angenommen, dass sich entlang des Links zwischen Sender und Empfänger der Nachricht eine Route etablieren lässt, obwohl dieser Link eine so hohe Fehlerrate aufweisen kann, dass es nicht sinnvoll ist, ihn zu verwenden. Ebenso kann ein Link aufgrund unterschiedlicher Sendeleistungen oder der Ausrichtung von Antennen, die nie perfekt omnidirektional abstrahlen, unidirektional sein, so dass der Pfad in der Rückrichtung nicht aufgebaut werden kann.

Um die Existenz von Links festzustellen und deren Qualität zu bestimmen, kann zum einen die Nutzlast selbst dienlich sein. Datenrahmen können schließlich von anderen Geräten in der Nachbarschaft des Senders auch dann empfangen werden, wenn sie nicht die beabsichtigten Empfänger sind, sofern alle Geräte einen gemeinsamen Übertragungskanal verwenden. Damit auch Verluste erkannt werden können, müssen alle Datenrahmen mit fortlaufenden Sequenznummern versehen werden. Dieser Ansatz ist beispielsweise für Sensornetze interessant, wenn davon ausgegangen werden kann, dass alle Sensoren ohnehin in regelmäßigen Abständen Daten erzeugen und versenden. Außerdem ist es in einer solch vergleichsweise homogenen Umgebung auch praktikabel, das Rahmenformat um ein Sequenznummernfeld zu erweitern.

Um hingegen in eher heterogenem Umfeld und unabhängig von der aktuellen Nutzlast die Qualität von Links bestimmen zu können, liegt es nahe, dass Geräte in regelmäßigen Zeitabständen eigens dafür vorgesehene Nachrichten versenden, die als Broadcast-Rahmen an alle Geräte adressiert sind, die in der Lage sind, sie zu empfangen. In der Literatur werden solche Nachrichten meist als *Beacons* oder *Hello-Messages* bezeichnet.

Die vermutlich erste Veröffentlichung zu diesem Thema ist [DRWT97]. Dort wird vorgeschlagen, dass die Geräte die Empfangssignalstärke der Beacons jedes Nachbarn durch eine Alterungsfunktion glätten und solche Links als „starke“ Links betrachten, deren zugehörige geglättete Signalstärke über einem bestimmten Schwellwert liegt. Links, entlang derer zwar Beacons empfangen werden, die aber nicht „stark“ in diesem Sinn sind, werden als „schwache“ Links betrachtet. Diese Kategorisierung nutzt das ebenfalls in [DRWT97] vorgeschlagene Routingprotokoll. Dabei wird aber außer Acht gelassen, dass Links aus verschiedenen Gründen asymmetrisch sein können.

Dies ist im Design des Mobile-Mesh-Link-Discovery-Protocol (MMLDP) [Gra00] berücksichtigt worden: Hier wird ab dem Zeitpunkt des Empfangs eines Beacons für eine vordefinierte Zeitspanne ein unidirektionaler Link angenommen. Diese Zeitspanne sollte länger als ein Beacon-Intervall sein, um eine stabile Sicht der Topologie zu ermöglichen. Damit Geräte Kenntnis über bidirektionale Links erlangen, beinhalten Beacons die Adressen jener Geräte, von denen unidirektionale Links zum Sender des Beacons bestehen. Damit sieht ein Ge-

rät einen Link als bidirektional an, wenn es über diesen Link Beacons empfängt, die seine eigene Adresse enthalten. Dieser Mechanismus, der auch in [GdWMJ03] und [GdWMJ05] verwendet wird, basiert offensichtlich auf der vereinfachten Annahme binärer Linkzustände und bewirkt damit, dass auch Links mit hohen Fehlerraten genutzt werden, die nur sporadisch fehlerfreie Übertragungen ermöglichen.

Um das zu verhindern, könnten wiederum Signalstärkemessungen in den Mechanismus integriert werden. Der Nachteil davon ist aber, dass einfache, preiswerte Hardware nicht unbedingt dazu in der Lage ist, solche Messungen überhaupt oder mit brauchbarer Genauigkeit durchzuführen. Ein anderer Ansatz ist die Abschätzung der Paketfehlerrate eines Links, wie sie sich auch im Bereich klassischer Festnetzprotokolle findet: Beim Internet Gateway Routing Protokoll (IGRP, [Rut91]) der Firma Cisco geht die Paketfehlerrate eines Links in die Routingmetrik ein. Gemessen wird sie als gleitender Durchschnitt der beobachteten Paketfehlerrate über feste Intervalle. In [WC03], [WTC03] wird vorgeschlagen, diese Methode für Sensornetze zu verwenden. Es werden verschiedene Schätzer untersucht, die teilweise mit geringerem Speicherbedarf auskommen als ein gleitender Durchschnitt, der sich jedoch ebenfalls als guter Schätzer erweist.

In [CE04] wird ein Mechanismus zur Linkerkennung vorgeschlagen, der auf einem dynamischen Schwellwert der gemessenen Paketverlustraten beruht: Ein Link gilt dann als zuverlässig, wenn ein Anteil von mindestens $1 - 1/N$ aller Pakete korrekt empfangen worden ist, wobei N die Anzahl der Nachbarn angibt. Je mehr Nachbarn ein Gerät also hat, desto geringer darf die Paketfehlerrate eines Links sein. Wie in Abschnitt 8.3.1.1 erläutert wird, hat ein solcher Mechanismus den Vorteil, dass der Schwellwert sich an die Gerätedichte anpasst: In Bereichen mit geringerer Gerätedichte werden genügend Links etabliert, um eine gute Konnektivität des Netzes sicherzustellen, während in Bereichen mit höherer Gerätedichte, wo eine Gerät also mehr potenzielle Nachbarn zur Auswahl hat, nur Links mit geringen Paketfehlerraten verwendet werden.

Anstatt die Eigenschaften von Links proaktiv zu erkunden, wird in [SNK05] vorgeschlagen, dies stärker im reaktiven Routing zu integrieren. Beispielsweise müsste das AODV-Protokoll (siehe Abschnitt 2.3.4) dann so modifiziert werden, dass ein Gerät eine RREQ-Nachricht nicht nur einmal, sondern mehrmals weiterleitet. Dadurch wird der durch das Beaconing verursachte Overhead eingespart, allerdings zum Preis einer noch höheren Netzlast bei der Routensuche. Der entscheidende Nachteil des in [SNK05] vorgestellten Ansatzes kann aber auch darin gesehen werden, dass die durch Beaconing gesammelte Information zur Netztopologie nicht zwingend nur von Routingmechanismen benötigt wird, sondern wie in der vorliegenden Arbeit auch von Mechanismen zur Sendeleistungsregelung.

2.3 Routing in drahtlosen Multihop-Netzen

Bis heute zeichnet sich drahtlose Kommunikation durch eine starke Abhängigkeit von Infrastruktur aus, deren Installation grundsätzlich mit Aufwand und mit Kosten verbunden ist.

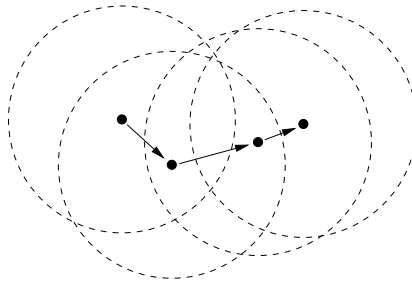


Abbildung 2.5: Beispiel eines drahtlosen Multihop-Netzes

Drahtlose Kommunikation ohne Verwendung zusätzlicher Infrastruktur ist im heutigen Alltagsgebrauch aber nur zwischen Geräten üblich, zwischen denen eine unmittelbare Funkverbindung besteht. Ein drahtloses Multihop-Netz zeichnet sich hingegen dadurch aus, dass die beteiligten Geräte ohne Verwendung einer Infrastruktur drahtlos miteinander kommunizieren, wobei jedes Gerät unter Umständen auch als Router fungiert, so dass zur Datenübertragung zwischen zwei Geräten kein direkter Funkkontakt nötig ist (siehe Abbildung 2.5).

2.3.1 Netzweites Fluten

Ein einfacher Mechanismus, um ein Paket an alle Netzteilnehmer zu senden, ist das *Fluten* (engl. *flooding*). Jedes Gerät leitet ein Paket bei diesem Vorgang genau einmal weiter, und zwar nach dem ersten Empfang. Es handelt sich hierbei um die einfachste Form des Routings, da sich so theoretisch auch Unicast-Pakete an ihr Ziel bringen ließen. Allerdings geht das Fluten mit einem sehr hohen Ressourcenverbrauch einher und ist daher höchstens in Szenarien angebracht, deren Dynamik so hoch ist, dass es kaum möglich ist, stabile Routen aufzubauen. Das Fluten wird allerdings von vielen Routingprotokollen gezielt zur Verteilung von Routinginformationen eingesetzt.

Effiziente Implementierungen des Flutens arbeiten üblicherweise so, dass die Geräte für einen begrenzten Zeitraum speichern, welche Pakete sie bereits bearbeitet haben, wobei Pakete häufig durch die Kombination aus Absender-Adresse und einer vom Absender vergebenen Sequenznummer identifiziert werden.

2.3.2 Proaktives und reaktives Routing

Klassische Routingprotokolle zeichnen sich dadurch aus, dass jedem Gerät grundsätzlich Routen zu allen anderen Geräten im Netz bekannt sind, sofern das System sich nicht gerade in einem Fehlerzustand befindet. Routingprotokolle mit dieser Eigenschaft werden als *proaktiv* bezeichnet, da den Geräten Routinginformationen bereits zur Verfügung stehen, bevor der Bedarf dazu besteht.

In der Forschung zu drahtlosen Multihop-Netzen ist das Paradigma des *reaktiven* (oder auch

On-demand-) Routings entstanden. Reaktive Routingprotokolle erhalten nur die Routen aufrecht, die tatsächlich für die Auslieferung von Datenpaketen benutzt werden. Dies hat den Vorteil, dass kein Aufwand zur Erhaltung von Routen betrieben wird, die gar nicht verwendet werden. Damit zahlt sich der Einsatz eines solchen Protokolls nur dann aus, wenn eine genügend hohe Zahl von Gerätepaaren keine Daten untereinander austauscht.

Proaktives Routing wird für drahtlose Multihop-Netze insbesondere deshalb häufig als ungeeignet erachtet, weil in diesen Netzen keine Hierarchien existieren, die es ermöglichen, Routen zu unterschiedlichen Zielen zusammenzufassen. Deshalb ist für jedes einzelne Ziel ein eigener Routingtabelleneintrag nötig, so dass der Aufwand für proaktives Routing sehr viel höher ist als in Netzen, die einen hierarchischen Adressraum verwenden, der die Netzstruktur widerspiegelt.

Einige Arbeiten in diesem Bereich beschäftigen sich zwar mit Verfahren, solche Hierarchien zu etablieren (ein guter Überblick findet sich in [SM02]), diese Verfahren sind jedoch recht kompliziert und ihre Anwendung in dynamischer Umgebung so aufwändig, dass ernsthaft bezweifelt werden darf, dass die Kosten in einem sinnvollen Verhältnis zum Nutzen stehen. Dies liegt daran, dass es in drahtlosen Multihop-Netzen im Allgemeinen kaum Strukturen gibt, die so stabil sind, dass sich darauf langlebige Hierarchien aufsetzen lassen. Dies kann höchstens in speziellen Szenarien anders sein, z.B. ist davon auszugehen, dass die Struktur militärischer Netze zu einem hohen Grad mit der hierarchischen Organisationsstruktur korrespondiert.

2.3.3 Link-State- und Distance-Vector-Protokolle

Bis auf einige exotischere Protokolle lassen sich die meisten Routingprotokolle in zwei Kategorien einteilen: *Link-State-Protokolle* basieren auf der netzweiten Verbreitung von Informationen über alle oder über ausgewählte Links, so dass Geräte für ihre Routingentscheidungen auf eine mehr oder weniger detaillierte Kenntnis der Netztopologie zurückgreifen können. Bei *Distance-Vector-Protokollen* hingegen tauschen benachbarte Geräte untereinander lediglich Distanzvektoren aus, die die *Kosten* von einer Quelle zu den verschiedenen Zielen enthalten. Diese werden meist in *Hops* gemessen, d.h. der Anzahl der Übertragungen entlang des Pfads zum Ziel. Distance-Vector-Protokolle verursachen damit weniger Overhead, müssen dafür aber auch mit einer sehr viel eingeschränkteren Sicht der Netztopologie auskommen. Ein Gerät errechnet die Kosten zu einem Ziel über einen bestimmten Nachbarn als Summe der Kosten, die der Nachbar für dieses Ziel angibt, und der Kosten des Links zum Nachbarn. Werden Distanzen in Hops gemessen, werden also alle von den Nachbarn empfangenen Kosten um 1 erhöht. Als ausgehender Link für eine Ziel wird dann der Link zu dem Nachbarn gewählt, über den das Ziel mit den geringsten Gesamtkosten erreichbar ist.

Distance-Vector-Protokolle haben den Nachteil, dass Routing-Zyklen entstehen können, die für lange Zeitspannen persistent sind. Dieses *Count-To-Infinity-Problem* entsteht dadurch, dass ungültige Routinginformationen prinzipiell beliebig lange von Gerät zu Gerät weitergegeben werden können. Eine genaue Beschreibung findet sich z.B. in [Tan03]. Lange Zeit galt dieses Problem als inhärent mit dem Distance-Vector-Ansatz verknüpft, und es existierten

lediglich Erweiterungen, um dieses Problem abzuschwächen, die es jedoch nicht vollständig lösen konnten. Interessanterweise fand sich gerade im Rahmen der Forschung zu drahtlosen Multihop-Netzen eine elegante und effektive Lösung [PB94].

Dieser Sequenznummern-Mechanismus erfordert, dass zu jedem Ziel nicht nur die Distanz, sondern auch eine *Ziel-Sequenznummer* ausgetauscht wird, die ein Maß für die Aktualität der Routinginformation ist. Damit erhöht sich der Routing-Overhead lediglich um einen konstanten Faktor.

Diese Sequenznummer wird vom Zielgerät selbst für jedes Routingpaket inkrementiert, und alle Geräte verwenden grundsätzlich nur die aktuellste ihnen bekannte Routinginformation. Geht ein Link zu einem Nachbarn verloren, werden die Distanzen zu allen Zielen, für die dieser Link als Ausgangslink diente, auf unendlich gesetzt und die zugehörigen Sequenznummern inkrementiert. Damit ist sichergestellt, dass keine Routinginformationen für diese Routen verwendet werden, die auf dem nicht mehr existenten Link beruhen.

2.3.4 Das AODV-Protokoll

Wie der Name des *Ad-hoc-On-demand-Distance-Vector-Protokolls* (AODV) bereits sagt, handelt es sich um ein reaktives Distance-Vector-Protokoll. Es setzt den oben beschriebenen Sequenznummern-Mechanismus ein, um das Count-To-Infinity-Problem zu vermeiden.

Liegt einer Quelle ein Paket für eine Ziel vor, zu dem gar kein oder ein ungültiger Routingtabelleneintrag vorhanden ist, versendet sie eine *Route-Request-Nachricht* (RREQ) als Broadcast-Rahmen an ihre Nachbarn. Die Nachricht enthält u.a. die Adressen der Quelle und des Ziels und eine minimale Zielsequenznummer, falls bereits eine Route und damit eine Sequenznummer für das Ziel bekannt waren. Damit die Zwischenstationen die *Reverse-Route* zurück zur Quelle aufbauen können, enthält die Nachricht auch die unmittelbar zuvor inkrementierte Sequenznummer der Quelle.

Eine empfangene RREQ-Nachricht wird von einem Gerät nur dann bearbeitet, wenn das Gerät nicht bereits an der gleichen Routensuche beteiligt gewesen ist, die eindeutig durch die Adresse der Quelle und einer in der Nachricht ebenfalls enthaltenen RREQ-ID identifiziert wird. Zunächst wird dann die Reverse-Route zur Quelle in die Routingtabelle eingetragen, wobei der Link zum Sender der Nachricht als Ausgangslink gespeichert wird. Ist der Empfänger selbst das Ziel oder ist ihm eine hinreichend aktuelle Route zum Ziel bekannt (deren Sequenznummer also nicht geringer als die in der RREQ-Nachricht angegebene Zielsequenznummer ist), so antwortet er mit einer *Route-Reply-Nachricht* (RREP), ansonsten wird die RREQ-Nachricht als Broadcast-Rahmen weitergeleitet.

Eine RREP-Nachricht wird als Unicast-Übertragung direkt an die in der RREQ-Nachricht angegebene Quelle entlang der bereits etablierten Reverse-Route verschickt. Alle Geräte entlang dieses Pfads aktualisieren ihren Routingtabelleneintrag für das Ziel mit der in der RREP-Nachricht enthaltenen Zielsequenznummer und dem Link zu demjenigen Gerät, von dem sie die RREP-Nachricht empfangen haben, als Ausgangslink. Außerdem wird die Adresse des

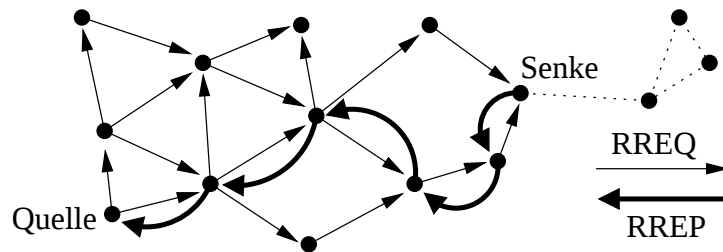


Abbildung 2.6: Veranschaulichung der Etablierung einer Route mit AODV

Geräts, an das die RREP-Nachricht weitergeleitet wird, in eine mit dem Ziel assoziierten *Precursor-List* eingefügt.

Schlägt das Versenden eines Pakets fehl, so geht das AODV-Protokoll davon aus, dass der entsprechende Link nicht mehr vorhanden ist, z.B. weil sich die beiden Geräte voneinander entfernt haben oder weil das Zielgerät ausgeschaltet wurde. Wurde dieser Link auch zur Weiterleitung der Pakete anderer Sender genutzt, so wird eine *Route-Error-Nachricht* (RERR) verschickt, die alle betroffenen Ziele und die dazugehörigen Sequenznummern enthält. Empfänger dieser Nachricht sind all jene Geräte, die in der *Precursor-List* eines betroffenen Ziels enthalten sind. Handelt es sich um genau ein Gerät, empfiehlt die AODV-Spezifikation eine gesicherte Unicast-Übertragung, sonst sollte ein Broadcast-Rahmen verwendet werden, um unnötigen Overhead zu vermeiden. Beim Empfang der Nachricht werden wiederum alle betroffenen Geräte aus den *Precursor-Lists* mittels einer neuen RERR-Nachricht darüber informiert, dass die entsprechenden Routen ungültig geworden sind. So wird diese Information entlang der Reverse-Routes zurück zu den Quellen propagiert, die daraufhin neue Routen suchen können.

Das AODV-Protokoll beinhaltet weitere Optimierungen, die hier nur kurz erwähnt werden. Der *Expanding-Ring-Search* ist ein Mechanismus, der den Ressourcenverbrauch durch das Fluten von RREQ-Paketen verringern soll. Dazu wird bei einer Routensuche zunächst nur eine lokale Umgebung (der sog. *Scope*) geflutet, indem das Time-To-Live-Feld (TTL) im IP-Header des RREQ-Pakets entsprechend gesetzt wird. Nach einem erfolglosen Suchvorgang wird dieser Scope vergrößert und eine neue Suche ausgelöst. Kann nach mehreren Wiederholungen dieses Vorgangs keine Route aufgebaut werden, wird schließlich doch das gesamte Netz geflutet.

Der optionale *Local-Repair* sieht vor, dass eine Zwischenstation auf einer Teilroute zu einem Ziel, die bemerkt, dass ein Link weggefallen ist, selbstständig die Route neu aufbaut, anstatt die Quellen mittels des RERR-Mechanismus darüber zu informieren, dass die Route nicht mehr gültig ist. Dies ist vor allem dann sinnvoll, wenn die Zwischenstation nur wenige Hops vom Ziel entfernt ist.

2.4 Topologien und deren Eigenschaften

2.4.1 Modellierung von Stationen, Links und Interferenz

Die Topologie eines Netzes kann als gerichteter Graph $G = (S, L)$ modelliert werden, wobei S die Menge der Stationen repräsentiert und $(v, w) \in L$ ist, wenn sich Station w innerhalb der Sendereichweite von Station v befindet, d.h. wenn w ein von v ausgesendetes Signal mit hoher Wahrscheinlichkeit fehlerfrei dekodieren kann, sofern dies nicht durch Interferenz anderer Stationen verhindert wird (siehe unten).

Ein *bidirektionaler Link* zwischen den Stationen v und w besteht dann, wenn sowohl $(v, w) \in L$ als auch $(w, v) \in L$ sind. Ist hingegen $(v, w) \in L$, aber $(w, v) \notin L$, so bezeichnet man (v, w) als einen *unidirektionalen Link* von v nach w . Damit ist zunächst keine Aussage darüber verbunden, ob unidirektionale Links genutzt werden können. Da Quittierungsmechanismen auf der Data-Link-Ebene nur entlang bidirektionaler Links möglich sind, ist die Annahme, dass nutzbare Links bidirektional sein müssen, aber relativ häufig anzutreffen. Auch der vorliegenden Arbeit liegt diese Annahme zugrunde. Lediglich die Tatsache, dass beim Fluten auch unidirektionale Links genutzt werden, weil es nicht unbedingt sinnvoll ist, Broadcast-Rahmen explizit zu quittieren, stellt eine Ausnahme dar.

Eine gängige Art der Modellierung von Netztopologien besteht darin, die Links aus den zuvor festgelegten Positionen in einem euklidischen Raum und Sendereichweiten der Stationen zu berechnen. Bezeichnet r_v die Sendereichweite der Station $v \in S$ und d_{vw} die Entfernung zwischen den Stationen $v, w \in S$, so ist die Menge der Links L gegeben durch

$$L = \{(v, w) \in S^2 \mid d_{vw} \leq r_v\}.$$

In einigen Publikationen wird Interferenz durch die vereinfachende Annahme modelliert, dass der Empfang einer Übertragung über einen Link $(v, w) \in L$ durch einen weiteren Sender $u \in S$ dann und nur dann gestört wird, wenn $(u, w) \in L$ ist. Ein Signal kann in der Realität jedoch über wesentlich größere Entfernungen störend wirken, als es mit hoher Wahrscheinlichkeit fehlerfrei dekodiert werden kann. An dieser Stelle wird meist wieder eine räumliche Betrachtungsweise verwendet, die (wie oben beschrieben) häufig schon zur Modellierung von Topologien durch gegebene Stationspositionen und -sendereichweiten eingesetzt wird.

In [GK00] werden zwei Interferenzmodelle eingeführt, die vielen Publikationen zugrunde liegen, das *physikalische Modell* und das *Protokollmodell*.

Das physikalische Modell lehnt sich stark an die physikalischen Anforderungen an, die zum fehlerfreien Empfang eines Signals erfüllt sein müssen. Eine Übertragung ist gemäß diesem Modell daher genau dann erfolgreich, wenn der SINR am Empfänger genügend hoch ist (siehe Abschnitt 2.1.2). Wenn also die Stationen $S' \subseteq S$ gleichzeitig senden, so kann das von $v \in S'$ ausgesendete Signal von $w \in S$ genau dann fehlerfrei empfangen werden, wenn

$$\frac{P_v/d_{vw}^\alpha}{N + \sum_{u \in S' \setminus \{v\}} (P_u/d_{uw}^\alpha)} \geq \beta,$$

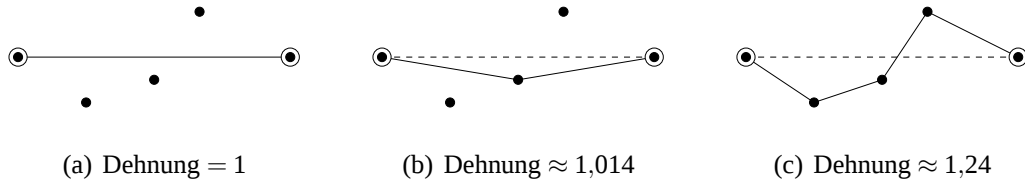


Abbildung 2.7: Veranschaulichung des Dehnungs-Begriffs

wobei P_u die Sendeleistung eines $u \in S'$, β die Untergrenze für den zum fehlerfreien Empfang benötigten SINR, N die Stärke des Hintergrundrauschens und α den Pathloss-Exponenten bezeichnen (siehe Abschnitt 2.1.3.1).

Im physikalischen Modell müssen also mehrere potenziell störende Sender gemeinsam betrachtet werden, so dass die Verwendung dieses Modells mit erheblichem Rechenaufwand verbunden sein kann. Im Protokollmodell werden Sender dagegen als störend oder nicht störend für eine Übertragung klassifiziert. Genau dann, wenn alle störenden Sender während einer Übertragung inaktiv sind, ist diese erfolgreich. In [GK00] ist das Protokollmodell zum einen für den Fall definiert, dass die Stationen für jede Übertragung eine optimale Sendeleistung wählen, und zum anderen für den Fall, dass alle Stationen eine feste Sendereichweite r haben. Im ersten Fall kann (unter den gleichen Voraussetzungen wie oben) eine Station w das Signal einer Station v fehlerfrei empfangen, wenn

$$d_{uw} \geq (1 + \delta) \cdot d_{vw}, \quad \forall u \in S' \setminus \{v\}.$$

Die Variable $\delta \geq 0$ ist die sog. *Guard-Zone*. Im zweiten Fall ist die Übertragung erfolgreich, wenn

$$d_{vw} \leq r \quad \text{und} \quad d_{uw} \geq (1 + \delta) \cdot r, \quad \forall u \in S' \setminus \{v\}.$$

Die Interferenzmodellierung unter der Annahme, dass Stationen für jede Übertragung eine günstige Sendeleistung wählen, ist nur mit zusätzlichen Einschränkungen sinnvoll. Da in diese Ungleichung nämlich nicht eingeht, dass der Empfänger der Übertragung eines $u \in S'$ sehr weit entfernt sein kann, lassen sich einfache Beispiele konstruieren, in denen zwei Übertragungen gleichzeitig in diesem Modell stattfinden können, obwohl sich ein Empfänger in der Sendereichweite beider Sender befindet.

Die Definition des Protokollmodells unter der Annahme, dass die Stationen eine feste Sendereichweite r haben, entspricht eher einer intuitiven Sichtweise. Konkret wird hier modelliert, dass das Signal eines Senders mit Reichweite r andere Empfänger über die Distanz $(1 + \delta) \cdot r$ stört. Diese Version des Protokollmodells dient deshalb als Basis für das Interferenzmodell, das den in Kapitel 6 beschriebenen Auswertungen zugrunde liegt.

2.4.2 Der Dehnungsfaktor

Als Maß für den Overhead, der mit einer Netztopologie verbunden ist, weil Pakete unter Umständen nicht direkt von Quelle zu Senke übertragen werden können, sondern von anderen Stationen weitergeleitet werden müssen, dient in vielen Arbeiten der sog. *Dehnungsfaktor* (engl. *stretch factor*, siehe z.B. [BDEK02]). Dieser basiert auf dem Begriff der *Dehnung*. Die Dehnung eines Pfads ist definiert als der Quotient zwischen der Länge des Pfads (also der Summe der euklidischen Distanzen entlang der Kanten des Pfads) und der euklidischen Distanz zwischen den Endpunkten des Pfads. Abbildung 2.7 zeigt beispielhaft die Dehnung entlang unterschiedlicher Pfade zwischen einem Knotenpaar. Die Dehnung eines Knotenpaars ist definiert als die Dehnung entlang des kürzesten Pfads, der das Knotenpaar verbindet, und der Dehnungsfaktor eines Graphen bezeichnet schließlich die maximale Dehnung über alle Knotenpaare des Graphen.

Der Dehnungsfaktor hat allerdings zwei Eigenschaften, die seine Eignung für die vorliegende Arbeit einschränken. Zunächst geht in diesen Wert nur dasjenige Knotenpaar ein, das mit maximaler Dehnung verbunden ist, so dass diese Metrik insbesondere für Worst-Case-Abschätzungen geeignet ist, weniger jedoch für die Untersuchung der Kapazität des gesamten Netzes.

Außerdem lässt sich leicht nachvollziehen, dass der Dehnungsfaktor eines vollständig verbundenen Netzes 1 ist (dies ist die Dehnung jedes Knotenpaars, das in einer solchen Topologie ja durch eine Kante verbunden ist), und dass er mit jeder Kante, die aus der Topologie entfernt wird, höchstens gleich bleibt oder steigt (da die Wege nur länger werden können). Hier werden also nur die durch das Routing verursachten Umwege berücksichtigt, nicht aber der Overhead, der durch unnötig hohe Sendereichweiten verursacht wird: Die Brutto-Gesamtdistanz, über die ein Paket übertragen wird, bestimmt sich streng genommen schließlich nicht durch die euklidischen Distanzen zwischen den einzelnen Stationen entlang der Route, sondern durch die Summe der Sendereichweiten der beteiligten Stationen (zumindest dann, wenn keine zusätzliche Optimierung der Sendeleistung für jede einzelne Übertragung erfolgt, siehe Abschnitt 3.2).

Den beiden genannten Aspekten trägt die in Abschnitt 6.3.2 eingeführte Metrik Rechnung.

2.4.3 Übertragungs-Scheduling

Für die in dieser Arbeit betrachteten Single-Channel-Netze müssen Übertragungen, die aufgrund von Interferenz nicht gleichzeitig erfolgreich stattfinden können, zeitlich getrennt werden. Diese Art der gemeinsamen Nutzung des Übertragungsmediums wird als *Time-Division-Multiple-Access* (TDMA) bezeichnet [NK85]. Ein sich periodisch wiederholendes *TDMA-Schema* oder *Übertragungs-Schedule* legt fest, in welchen Zeiträumen bzw. *Zeitslots* welche Datenübertragungen stattfinden können. Jede Zeitspanne, in der das gesamte TDMA-Schema einmal abgearbeitet wird, wird *TDMA-Frame* genannt. Abbildung 2.8 zeigt beispielhaft ein solches TDMA-Schema, das zum noch zu erläuternden Beispiel in Abbildung 2.9(b) korre-

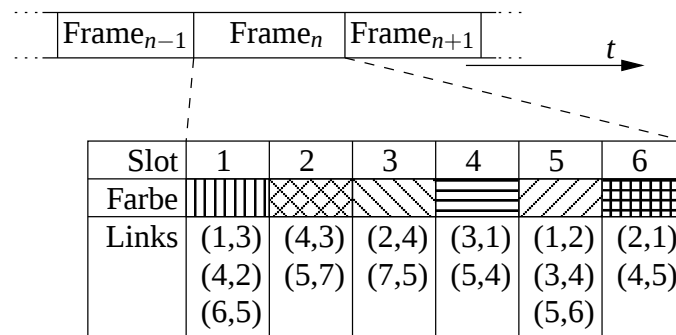


Abbildung 2.8: Beispiel eines sich periodisch wiederholenden TDMA-Schemas

spondiert. In der Praxis müssen die einzelnen Zeitslots durch Schutzzeiten zu Beginn und zum Ende eines Slots voneinander getrennt werden, um trotz der Ungenauigkeiten im Timing einzelner Stationen und gegebener Signallaufzeiten Störungen zu verhindern. Zum Vergleich des möglichen Durchsatzes in verschiedenen Topologien müssen feste Schutzzeiten aber nicht berücksichtigt werden, da sie den Durchsatz nur um einen konstanten Faktor verringern.

Einem Link steht langfristig der Anteil der Linkbandbreite zur Verfügung, der dem Zeitan- teil dieses Links im TDMA-Schema entspricht. Anders ausgedrückt ist die einem Link zur Verfügung stehende Bandbreite durch die Datenmenge bestimmt, die über die Gesamtdauer des TDMA-Schemas übertragen werden kann. Daher ist eine gegebene Last, d.h. eine Menge von Links mit bestimmten Bandbreitenanforderungen, theoretisch dann realisierbar, wenn ein gültiges TDMA-Schema existiert, das auf jedem Link die geforderte Bandbreite zulässt.

Die explizite Verwendung eines TDMA-Schemas erfordert eine Synchronisation aller Stationen, die nur schwierig umzusetzen ist, falls die Stationen über einen Bereich verteilt sind, der groß ist im Verhältnis zur Sendereichweite eines Geräts. Ein bestimmtes TDMA-Schema ist außerdem nur für eine Teilmenge der realisierbaren Lasten geeignet; ändert sich also die Last, muss unter Umständen also ein neues TDMA-Schema erzeugt werden. TDMA-Schemata dienen in dieser Arbeit daher nur der Untersuchung des Durchsatzes, der sich in einer Topologie im günstigsten Fall realisieren lässt (siehe Kapitel 5 und 6).

Unter der Annahme, dass alle Übertragungen die gleiche Zeitspanne benötigen, ist ein TDMA-Schema formal eine Zuordnung $s : T = \{1, \dots, m\} \rightarrow 2^L$ von m Zeitslots fester Länge zu gleichzeitig stattfindenden Übertragungen (2^L bezeichnet die Potenzmenge von L , der Menge aller Links). Algorithmen zur Erzeugung von Schedulings werden häufig als Graphfärbelgorithmen formuliert. Farben korrespondieren dabei mit Zeitslots; bei einer Kantenfärbung des Topologiegraphen bedeutet das beispielsweise, dass gleichzeitige Übertragungen genau auf den Links stattfinden, die gleich gefärbt sind. Solche Färb- bzw. Schedulingprobleme sind im Allgemeinen NP-hart [Ari84], so dass in diesem Zusammenhang meist Approximationsalgorithmen verwendet werden.

Um eine möglichst hohe Auslastung zu erreichen, muss eine gültige Färbung mit einer mög-

lichst geringen Gesamtzahl von Farben berechnet werden. In den Färbungen, die in diesem Abschnitt vorgestellt werden, wird jeder Entität genau eine Farbe zugeordnet, es gibt aber auch Algorithmen, die einer Entität mehrere Farben zuordnen (beispielsweise wird in [SC99] die Einschränkung der strikten Fairness gelockert und auf diese Weise ein höherer Gesamtdurchsatz erzielt).

Ein TDMA-Schema ist genau dann *gültig*,

1. wenn Übertragungen nur zwischen Stationen vorgesehen sind, zwischen denen ein Link besteht und
2. wenn gleichzeitige Übertragungen nur dann vorgesehen sind, wenn sie sich nicht gegenseitig stören und wenn die drahtlose Hardware dies zulässt:

- a) Geräte sind nicht in der Lage, gleichzeitig zu senden und zu empfangen, d.h.

$$\forall t \in T, \forall (v, w), (v', w') \in s(t) : w \neq v'.$$

Eine gültige Kantenfärbung des Topologiegraphen muss also für eine ein- und eine ausgehende Kante an einem Knoten unterschiedliche Farben verwenden.

- b) Geräte sind nicht in der Lage, von mehreren Sendern gleichzeitig erfolgreich zu empfangen, d.h.

$$\forall t \in T, \forall (v, w), (v', w') \in s(t) : w \neq w'.$$

Eine gültige Kantenfärbung des Topologiegraphen muss also für eingehende Kanten an einem Knoten unterschiedliche Farben verwenden.

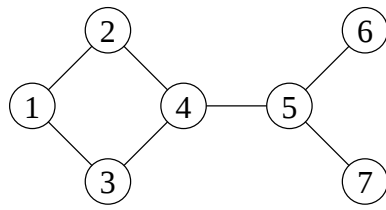
- c) Bei Unicast-Übertragungen kann ein Sender zu einem bestimmten Zeitpunkt nur an einen Empfänger übertragen, d.h.

$$\forall t \in T, \forall (v, w), (v', w') \in s(t) : v \neq v'.$$

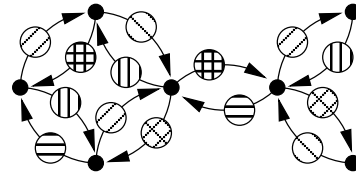
Eine gültige Kantenfärbung des Topologiegraphen muss in diesem Fall also für ausgehende Kanten an einem Knoten unterschiedliche Farben verwenden.

Im Fall von Unicast-Übertragungen müssen adjazente Kanten also generell unterschiedlich gefärbt sein. Eine solche Färbung wird als *(Distance-1)-Edge-Colouring* und die dafür mindestens benötigte Anzahl an Farben als *chromatischer Index* des Graphen bezeichnet. Abbildung 2.9(b) zeigt beispielhaft eine solche Kantenfärbung. Da der höchste Knotengrad eine Untergrenze des chromatischen Index ist, ist die Anzahl unterschiedlicher Farben in diesem Beispiel minimal.

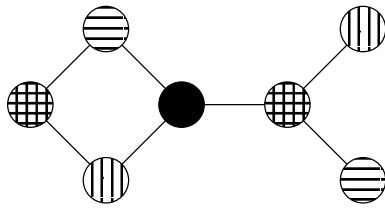
Im Broadcast-Fall ist eine Kantenfärbung unter Berücksichtigung der ersten beiden Einschränkungen denkbar. Adjazente Kanten müssen demnach unterschiedlich gefärbt sein, es sei denn, beides sind ausgehende Kanten des gemeinsamen Knotens. Zur Verringerung des Aufwands bietet sich aber eine Knotenfärbung an. Eine Farbe identifiziert dann einen Zeitslot,



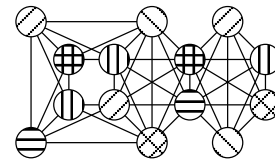
(a) Netztologie



(b) Distance-1-Edge-Colouring



(c) Distance-2-Vertex-Colouring



(d) Distance-1-Vertex-Colouring eines zugehörigen Konfliktgraphen

Abbildung 2.9: Veranschaulichung verschiedener Graphfärbungen

in dem bestimmte Stationen einen Broadcast-Rahmen an ihre Nachbarn senden. Wegen der ersten beiden Einschränkungen ist klar, dass adjazente Knoten unterschiedliche Farben haben müssen, ebenso wie Knoten $v, w \in S$, für die ein $u \in S$ existiert mit $(v, u), (w, u) \in L$. Eine gültige Knotenfärbung lässt sich in eine gültige Kantenfärbung überführen, indem jede Kante die Farbe des Knotens erhält, von dem sie ausgeht. Im Spezialfall, dass alle Links bidirektional sind, muss jeder kürzeste Pfad zwischen Knoten gleicher Färbung mindestens zwei weitere Knoten anderer Färbung enthalten, weshalb man hier vom *Distance-2-Vertex-Colouring* spricht. Eine solche Färbung ist in Abbildung 2.9(c) beispielhaft dargestellt. Dadurch, dass weniger Einschränkungen beachtet werden müssen als im Fall von Unicast-Übertragungen, ist die Zahl der benötigten Farben im Allgemeinen geringer.

Für verwandte Probleme, z.B. der Frequenzzuweisung in zellulären Netzen, können ähnliche Beschränkungen für Graphfärbungen formuliert werden. Die Arbeit [Ram97] stellt einen allgemeinen Algorithmus zur Berechnung von Färbungen vor, der verschiedene Annahmen umsetzen kann.

Die oben genannten Färbgorithmen modellieren ausschließlich die Einschränkungen der drahtlosen Hardware, eine zusätzliche Interferenzmodellierung ist darin aber nicht enthalten. Das Distance-1-Edge-Colouring zum Scheduling von Unicast-Übertragungen ist deshalb lediglich für Multichannel-Netze geeignet, in denen gleichzeitige Übertragungen sich nicht gegenseitig stören. In Single-Channel-Netzen ist es dagegen sehr wahrscheinlich, dass die Signale mehrerer Sender, die zu einem Empfänger adjazent sind, sich bei diesem so überlagern, dass das für diese Station bestimmte Signal nicht erfolgreich dekodiert werden kann. In Abbildung 2.9(b) beispielsweise sind die Links (3,4) und (5,6) gleichzeitig aktiv; da Station 4

auch in Sendereichweite von Station 5 liegt, ist es je nach konkreter Dämpfung entlang der Links und dem von der Hardware tolerierbaren SINR möglich, dass das Signal von Station 5 den Empfang an Station 4 stört.

Für eine generische Interferenzmodellierung bietet sich die Darstellung der Einschränkungen für gleichzeitige Übertragungen in einem *Konfliktgraphen* an, dessen Knotenmenge die zu färbenden Entitäten des Topologiegraphen sind. Eine Kante zwischen zwei Knoten besteht dann und nur dann, wenn diese Entitäten nicht zur gleichen Zeit aktiv sein dürfen. Ein Distance-1-Vertex-Colouring des Konfliktgraphen repräsentiert dadurch ein gültiges Scheduling (siehe [SC99], [WYG⁺01], [JPPQ03]). Die Mindestzahl der dafür benötigten Farben wird als *chromatische Zahl* des Konfliktgraphen bezeichnet.

Geht es um das Scheduling von Unicast-Übertragungen, ist der Konfliktgraph gegeben durch $G_C = (L, E_C)$. Abbildung 2.9(d) zeigt den zur Topologie aus Abbildung 2.9(a) korrespondierenden Konfliktgraphen für den Fall, dass zwei Links sich genau dann gegenseitig stören, wenn sie adjazent sind. Aus diesem Grund korrespondiert die Knotenfärbung dieses Konfliktgraphen zur Kantenfärbung des Topologiegraphen (siehe Abbildung 2.9(b)), eine andere Interferenzmodellierung ist jedoch leicht durch eine geeignete Erweiterung der Kantenmenge des Konfliktgraphen möglich.

Generell können auf diese Art und Weise beliebige Einschränkungen modelliert werden, die die folgende Bedingung erfüllen: Ob zwei Entitäten gleichzeitig aktiv sein können oder nicht, muss unabhängig davon sein, welche anderen Entitäten noch aktiv sind. So lässt sich beispielsweise das Protokollmodell mit Konfliktgraphen implementieren, das physikalische Modell aber nicht (siehe Abschnitt 2.4.1).

Es ist nicht verwunderlich, dass die Gesamtauslastung des Systems mit fallenden Sendeleistungen monoton steigt, da wegen des steigenden Spatial-Reuse immer mehr gleichzeitige Übertragungen stattfinden können. Im Broadcast-Fall bewirkt dies tatsächlich eine Steigerung des möglichen Durchsatzes, im Fall von Unicast-Strömen müssen jedoch auch die verlängerten Routen zwischen Sendern und Empfängern von Datenströmen berücksichtigt werden. Die Frage, ob oder wie viel von dem Gewinn an Spatial-Reuse sich in der Unicast-Kommunikation bemerkbar macht, wurde für gegebene Netztopologien bisher kaum erforscht. Die Schwierigkeit liegt darin, dass für gegebene Ströme im Allgemeinen eben nicht alle Stationen mit gleicher Häufigkeit senden oder Links mit gleicher Häufigkeit verwendet werden und dass die Rückkopplung zwischen dem Scheduling von Übertragungen und den Durchsätzen der Datenströme schwer in den Griff zu bekommen ist. Um diese Lücke zu schließen, wird in Abschnitt 5.3 ein neuer Algorithmus zur Erzeugung von TDMA-Schemata präsentiert, die den einzelnen Links in Abhängigkeit der vorliegenden Datenströme unterschiedliche Bandbreitene anteile zuordnen, wobei die Aufteilung der Ressourcen zwischen den Datenströmen (in der in Abschnitt 2.6 definierten Weise) fair ist.

2.4.4 Kapazität

Allgemein bezeichnet das Wort „Kapazität“ einen maximal erreichbaren Wert. Die Kapazität eines Netzes bezeichnet entsprechend den maximal erzielbaren Durchsatz. Dieser hängt entscheidend von der Art des Verkehrs ab. Da es grundsätzlich verschiedene Verkehrsarten gibt, werden in dieser Arbeit zwei verschiedene Kapazitätsbegriffe verwendet: Die *Broadcast-Kapazität* gibt den durch das Fluten von einem Absender zu allen anderen Netzteilnehmern maximal erzielbaren Durchsatz an, während unter der *Transportkapazität* der maximal erzielbare, aggregierte Ende-zu-Ende-Durchsatz von Unicast-Datenströmen verstanden wird.

Multicast-Verkehr, der an eine Teilmenge der Netzteilnehmer adressiert ist, stellt eine Mischform zwischen Unicast- und Broadcast-Verkehr dar. Sind die Multicast-Gruppen vergleichsweise klein, wird die Modellierung des Verkehrs durch mehrere Unicast-Ströme keine wesentlich anderen Ergebnisse liefern, was die Kapazität angeht. Sind die Multicast-Gruppen vergleichsweise groß, lässt sich der Verkehr entsprechend gut als Broadcast-Verkehr modellieren, da selbst die Stationen, die nicht Mitglieder einer Multicast-Gruppe sind, mit hoher Wahrscheinlichkeit am Transport des Multicast-Verkehrs beteiligt sind.

Zur Bestimmung der Netzkapazität in einem konkreten Szenario eignen sich grundsätzlich zwei Methoden: Zum einen kann der Datenverkehr in einem gegebenen Netz simuliert werden, d.h. es werden mit Hilfe einer Simulationssoftware Übertragung und Pufferung von Daten- und Kontrollpaketen unter Verwendung bestimmter Protokolle auf verschiedenen Ebenen des Schichtenmodells detailliert nachgebildet. Zum anderen kann die Kapazität eines Netzes mit Hilfe graphtheoretischer Ansätze errechnet oder approximiert werden (wie im vorangegangenen Abschnitt 2.4.3 beschrieben).

Die Messung der Netzkapazität mit Hilfe der Simulation gibt zwar möglicherweise einen guten Einblick in die Leistung eines Netzes mit einer bestimmten Topologie unter bestimmten Voraussetzungen, die Aussagen sind zunächst jedoch nur für die speziellen Annahmen gültig, die für die Simulation getroffen wurden. Wie in Abschnitt 2.2.1 beschrieben, sind IEEE 802.11 und das Transportprotokoll TCP für drahtlose Multihop-Netze nicht besonders gut geeignet, und viel Forschungsaufwand wird in die Verbesserung dieser Protokolle und in die Entwicklung alternativer Protokolle investiert ([XS01], [SHS04], [NHK05], [FLZ⁺05]). Deshalb sind Simulationsergebnisse, die auf diesen Protokollen basieren, für zukünftige drahtlose Multihop-Netze vielleicht nur wenig aussagekräftig.

Die Kapazität eines drahtlosen Netzes ergibt sich nach graphtheoretischen Ansätzen durch ein optimales Scheduling der Datenübertragungen, das möglichst viele Übertragungen gleichzeitig stattfinden lässt, die nicht miteinander interferieren (siehe Abschnitt 2.4.3). Ein Überblick über Arbeiten zu diesem Thema findet sich in Abschnitt 4.1.2.

An dieser Stelle muss noch erwähnt werden, dass auch Arbeiten existieren, die auf einem anderen Kapazitätsbegriff basieren. Während meist implizit angenommen wird, dass Datenpakete in erster Linie durch die Übertragung zwischen Stationen transportiert werden, wird in [GT01] festgestellt, dass die Kapazität eines drahtlosen Multihop-Netzes wesentlich größer ist, wenn die Mobilität der Geräte selbst genutzt wird, um Daten zu transportieren, wenn also

Datenpakete potenziell beliebig lange gepuffert werden, so dass sich die Topologie des Netzes durch die Mobilität der Geräte in dieser Zeit signifikant ändert. Dadurch werden die Paketverzögerungen beliebig groß. In [PB03] wird die Abhängigkeit zwischen maximaler Verzögerung und Kapazität genauer untersucht. Da diese Form der Datenübertragung nur für sehr spezielle Anwendungen geeignet ist, die sehr lange Verzögerungen in Kauf nehmen können, wird sie in der vorliegenden Arbeit nicht betrachtet.

2.5 Modellierung von Mobilität

Zur Bewertung von Protokollen in drahtlosen Multihop-Netzen oder zur Analyse von Eigenschaften solcher Netze ist es notwendig, die Mobilität der Geräte zu modellieren. Aufzeichnungen von Gerätebewegungen in konkreten Szenarien existieren mit dem nötigen Detailgrad bisher nicht. Die Betreiber zellulärer Netze verfügen über Daten zur Mobilität der Geräte in ihren Netzen, die aber lediglich auf der Ebene einzelner Mobilfunkzellen gegeben sind, so dass sich daraus kaum auf die Konnektivität solcher Geräte untereinander mittels einer WLAN-Technologie schließen lässt. Deshalb wird in der Forschung zu drahtlosen Multihop-Netzen auf Mobilitätsmodelle zurückgegriffen, mit denen sich Bewegungsdaten mit gewissen Eigenschaften generieren lassen, um die Leistungsfähigkeit von Protokollen in Simulationen zu überprüfen.

2.5.1 Das Random-Waypoint-Modell

Das Random-Waypoint-Modell wird in [JM96] eingeführt und ist ein sehr gängiges Mobilitätsmodell zur Generierung von Bewegungen zur Simulation drahtloser Multihop-Netze. Nach diesem Modell bewegt sich ein Gerät mit gleichbleibender Geschwindigkeit zu einem zufällig gewählten Zielpunkt auf der Simulationsfläche. Dort verweilt es für eine gewisse Zeit, bevor sich dieser Vorgang für einen neu gewählten Zielpunkt wiederholt.

Verschiedene Varianten des Random-Waypoint-Modells unterscheiden sich in einigen Details: Die gewählten Geschwindigkeiten ebenso wie die Verweildauern können verschiedenen Verteilungen entnommen oder auch festgelegt sein. Durch Veränderung beider Größen lässt sich die Dynamik des Netzes variieren, da sowohl höhere Geschwindigkeiten als auch kürzere Verweildauern die Dynamik erhöhen.

Das Random-Waypoint-Modell hat eine Eigenschaft, wegen der es für die vorliegende Arbeit ungeeignet ist. Die Geräte sind in eingeschwungenem Zustand nämlich nicht gleichmäßig über die Simulationsfläche verteilt, sondern die Gerätedichte ist zur Mitte der Simulationsfläche hin höher. Dies liegt informell darin begründet, dass Punkte, die der Mitte der Simulationsfläche näher sind, mit höherer Wahrscheinlichkeit auf einer Linie zwischen zwei Punkten liegen, die zufällig gemäß einer Gleichverteilung auf der Fläche gewählt werden. Dieses Phänomen ist in [BW02] untersucht worden, und es lässt sich sehr einfach und anschaulich für den eindimensionalen Fall zeigen: Die Wahrscheinlichkeit, dass ein Punkt x zwischen zwei

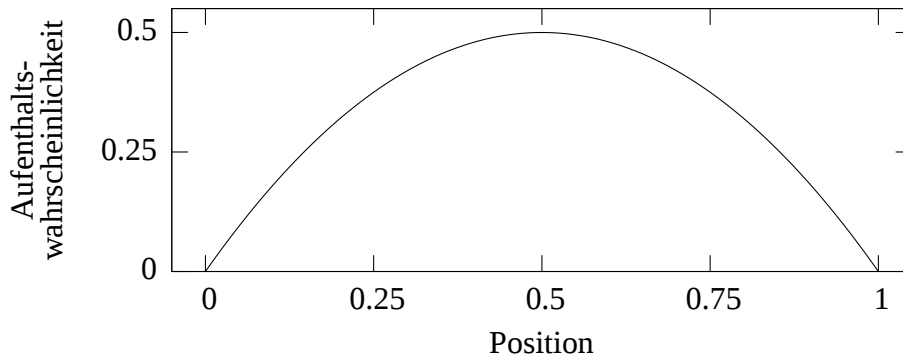


Abbildung 2.10: Aufenthaltswahrscheinlichkeiten im eindimensionalen Random-Waypoint-Modell (ohne Verweilen am Ziel)

zufällig gewählten $a, b \in [0; 1]$ liegt, ist gerade

$$P(a \leq x) \cdot P(b \geq x) + P(b \leq x) \cdot P(a \geq x) = 2 \cdot x \cdot (1 - x).$$

Der Verlauf dieser Funktion für $x \in [0; 1]$ ist in Abbildung 2.10 dargestellt.

Da die Zielpositionen, an denen die Geräte verweilen, über die Simulationsfläche gleichverteilt sind, ist die Verteilung der Geräte im eingeschwungenen Zustand eine Überlagerung zwischen der glockenförmigen Verteilung sich bewogender Geräte und der Gleichverteilung sich nicht bewogender Geräte. In welchem Verhältnis diese beiden Verteilungen in die resultierende Verteilung eingehen, hängt vom Verhältnis der durchschnittlichen Bewegungsdauer und der durchschnittlichen Verweildauer ab. Daher ist es ungünstig, eine variierende Dynamik durch verschiedene Werte für die Verweildauern der Geräte zu modellieren. Unter diesen Umständen variiert nämlich auch die Verteilung der Geräte, die ebenfalls einen Einfluss auf die Kapazität des Netzes hat. Je geringer die Aufenthaltswahrscheinlichkeit im Randbereich der Simulationsfläche ist, desto geringer ist beispielsweise der Erwartungswert für die Distanz zwischen einem Sender und einem Empfänger, so dass die Kommunikationspfade im Mittel weniger Zwischenstationen enthalten. Deshalb wird dieses Mobilitätsmodell für die vorliegende Arbeit nicht verwendet.

2.5.2 Das Gauß-Markov-Mobilitätsmodell

In [LH99] wird im Rahmen der Arbeit zu Mobilitätsmanagement in zellulären Netzen das Gauß-Markov-Mobilitätsmodell eingeführt. Mit geringfügigen Erweiterungen kann es aber ebenso verwendet werden, um die Bewegungen von Geräten auf einer vorgegebenen Fläche zu modellieren.

Gemäß diesem Modell wird der Bewegungsvektor eines Geräts in regelmäßigen Zeitabständen aktualisiert. Jede Dimension dieser Vektorenfolge ist charakterisiert durch einen speziel-

len autoregressiven Prozess erster Ordnung:

$$\begin{aligned} x_0 &= \mu + \sigma z_0, \\ x_n &= \alpha x_{n-1} + (1 - \alpha)\mu + \sqrt{1 - \alpha^2} \sigma z_n \quad \text{für } n > 0. \end{aligned}$$

Die z_i sind voneinander unabhängige, standardnormalverteilte Zufallszahlen (anders ausgedrückt handelt es sich bei $\{z_i : i \in \mathbb{N}_0\}$ um einen White-Noise-Prozess). Seinen Namen trägt das Gauß-Markov-Mobilitätsmodell deshalb, weil ein solcher diskreter Prozess die Abtastung eines kontinuierlichen Gauß-Markov- bzw. Ornstein-Uhlenbeck-Prozesses darstellt (wobei das Abtastintervall durch den Parameter α gesteuert wird) [Bar01]. Es ist $E(x_i) = \mu$ und $\text{Var}(x_i) = \sigma^2$. Außerdem hat der Prozess folgende Eigenschaften:

1. **Gauß-Eigenschaft:** Der Vektor $(x_{i_1}, \dots, x_{i_n})$ folgt für alle $i_1, \dots, i_n \in \mathbb{N}_0$ einer multivariaten Normalverteilung.
2. **Markov-Eigenschaft:** Die zukünftige Entwicklung des Prozesses hängt nur vom aktuellen Wert ab und nicht von der Entwicklung in der Vergangenheit, d.h.

$$P(x_{i_n} < c \mid x_{i_1}, \dots, x_{i_{n-1}}) = P(x_{i_n} < c \mid x_{i_{n-1}}), \forall i_1, \dots, i_n \in \mathbb{N}_0, i_1 < \dots < i_n.$$

3. **Stationarität:** Die Vektoren $(x_{i_1}, \dots, x_{i_n})$ und $(x_{i_1+h}, \dots, x_{i_n+h})$ sind für alle $h, i_1, \dots, i_n \in \mathbb{N}_0$ identisch verteilt.

Eine Folge von (m -dimensionalen) Bewegungsvektoren wird entsprechend erzeugt durch

$$\begin{aligned} v_0 &= \mu + \sigma w_0, \\ v_n &= \alpha v_{n-1} + (1 - \alpha)\mu + \sqrt{1 - \alpha^2} \sigma w_n \quad \text{für } n > 0, \end{aligned}$$

wobei μ hier den m -dimensionalen Erwartungswert der Folge bezeichnet und w_i Vektoren aus m unabhängigen, standardnormalverteilten Zufallswerten sind.

Die explizite Festlegung einer Höchstgeschwindigkeit ist in diesem Modell nicht vorgesehen und auch nicht nötig, da die Wahrscheinlichkeit für das Auftreten beliebig hoher Geschwindigkeiten vernachlässigbar gering ist. (Da die Komponenten der Bewegungsvektoren einer Normalverteilung mit Mittelwert μ und Varianz σ^2 entstammen, sind die Geschwindigkeiten Rayleigh-verteilt, falls $\mu = 0$, sonst sind sie Rice-verteilt. Siehe dazu auch die Ausführungen in Abschnitt 2.1.3.2.)

Für die Simulationen, die im Rahmen dieser Arbeit durchgeführt wurden, musste das Modell so erweitert werden, dass die Bewegungen der Geräte auf eine vorgegebene Simulationsfläche beschränkt bleiben. Dazu werden die Geräte zunächst gleichverteilt auf der Simulationsfläche platziert, und die Erzeugung der zweidimensionalen Bewegungsvektoren erfolgt wie oben beschrieben. Um zu verhindern, dass ein Gerät sich über den Rand der Simulationsfläche hinausbewegt, werden der aktuelle Bewegungsvektor v_n ebenso wie der Erwartungsvektor μ am Rand der Simulationsfläche gespiegelt (siehe Abbildung 2.11(a)). Ist die Simulationsfläche

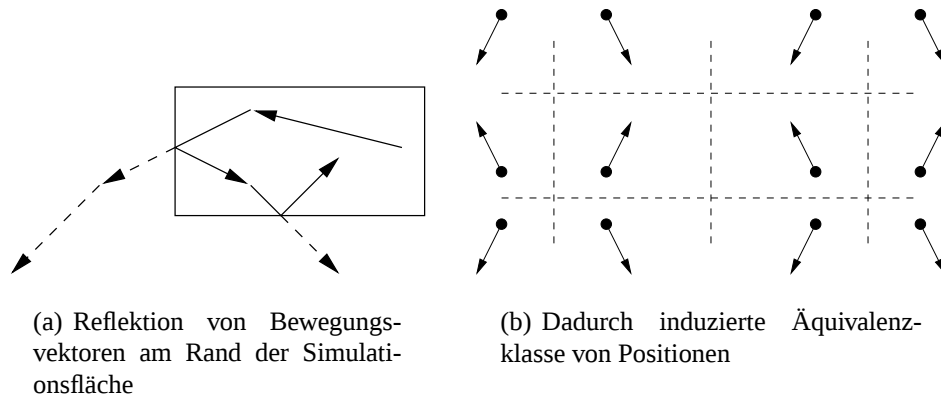


Abbildung 2.11: Einschränkung der Gerätepositionen auf vorgegebenes Rechteck

rechteckig, wird dies durch die Umkehrung des Vorzeichens einer bestimmten Komponente jedes der Vektoren erreicht.

Wie im Folgenden begründet wird, unterliegt die Verteilung der Gerätepositionen bei dieser Modellierung zu jeder Zeit einer Gleichverteilung, wenn die Phasen aller μ zu Beginn zufällig einer Gleichverteilung entnommen werden. Die beschriebene Vorgehensweise kann nämlich auch so interpretiert werden, dass die Position eines Geräts nicht durch einen Punkt auf einer endlichen Fläche $[0; X] \times [0; Y]$ gegeben ist, sondern durch eine Äquivalenzklasse von Punkten in $\mathbb{R}^2$, die genau ein Element auf der Simulationsfläche enthält. Dabei ist $(x, y) \equiv (x', y')$ genau dann, wenn

$$\exists i, j \in \mathbb{Z} : (x' = 2iX + x \vee x' = 2iX - x) \wedge (y' = 2jY + y \vee y' = 2jY - y).$$

Die Äquivalenzklasse einer Zielposition, die aus der Addition einer beliebigen Folge von Bewegungsvektoren auf eine Ursprungsposition resultiert, ist unabhängig davon, ob der oben beschriebene Spiegelungsmechanismus verwendet wird oder nicht. Abbildung 2.11(b) stellt eine solche Äquivalenzklasse und die sich durch Anwendung eines Bewegungsvektors auf ein Element dieser Klasse ergebende neue Äquivalenzklasse beispielhaft dar.

Wenn sowohl die Startpositionen auf der Simulationsfläche als auch die Phasen aller μ gleichverteilt sind, dann muss auch die Vereinigung der Äquivalenzklassen aller Positionen gleichverteilt in $\mathbb{R}^2$ sein. Diese Gleichverteilung der Gerätepositionen gilt damit auch für die Projektion auf die endliche Simulationsfläche.

Zur Modellierung der Mobilität für die Simulationsszenarien in der vorliegenden Arbeit wird daher für jedes Gerät ein eigenes, zufällig gewähltes μ verwendet, dessen Betrag allerdings ein Eingabeparameter und für alle Geräte gleich ist. Für $|\mu| = 0$ sind die Szenarien eher statisch in dem Sinn, dass die Geräte sich langfristig betrachtet nicht zielgerichtet bewegen. Je größer $|\mu|$ gewählt wird, desto zielgerichteter bewegen sich die Geräte in eine (für jedes Gerät individuelle) Richtung, bis sie am Rand der Simulationsfläche angelangt sind. Dies ist beispielhaft in Abbildung 2.12 veranschaulicht. Dort sind die Bewegungen zweier Geräte auf

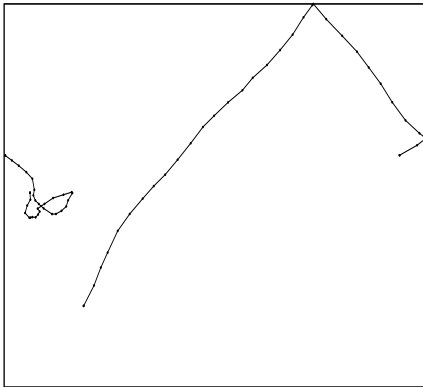


Abbildung 2.12: Beispiele für Bewegungen gemäß dem Gauß-Markov-Mobilitätsmodell

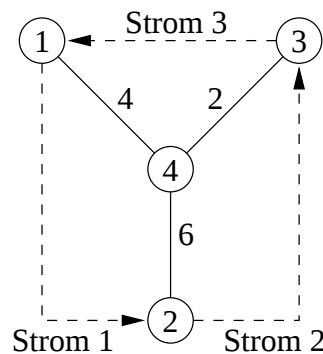


Abbildung 2.13: Beispiel zur Veranschaulichung von Fairness-Definitionen

einer quadratischen Simulationsfläche mit der Seitenlänge von 200m eingezeichnet. Die Bewegungsvektoren wurden alle 20s angepasst (in der Abbildung sind diese Stellen mit Punkten gekennzeichnet). In beiden Fällen sind $\alpha = 0,75$ und $\sigma = 0,1\text{m/s}$. Für das Gerät, das sich langfristig nicht aus einem relativ kleinen Gebiet herausbewegt, ist $|\mu| = 0$, während die zielgerichtete Bewegung mit $|\mu| = 0,5\text{m/s}$ erzeugt wurde.

2.6 Fairness

Ein Rechnernetz stellt Ressourcen zur Verfügung, die sich mehrere Instanzen (also Nutzer oder Prozesse eines Nutzers) teilen. Die Ressourcen können auf verschiedene Weisen auf die Instanzen aufgeteilt werden, wobei gewisse Einschränkungen beachtet werden müssen, damit eine Aufteilung auch gültig, also umsetzbar ist. Bei einem Netz aus Rechnern, die durch dedizierte Links miteinander verbunden sind, ist eine Aufteilung der Bandbreite zwischen verschiedenen Strömen mit vorgegebenen Routen dann gültig, wenn zum einen kein Strom mehr Ressourcen erhält, als er benötigt, und wenn zum anderen die Summe der Bandbreiten auf jedem Link die Gesamtkapazität des Links nicht übersteigt. Es ist wesentlich leichter zu erkennen, ob eine Aufteilung gemäß dieser Bedingungen gültig ist als für drahtlose Netze mit gemeinsam genutztem Übertragungsmedium, bei denen dies durch die Existenz eines gültigen Übertragungs-Schedules bestimmt ist (siehe Abschnitt 2.4.3). Deshalb dient ein kleines Rechnernetz mit dedizierten Links als Beispiel zur Erläuterung der verschiedenen Fairness-Begriffe.

In dem in Abbildung 2.13 dargestellten Beispielnetz soll die zur Verfügung stehende Bandbreite auf drei Ströme verteilt werden. Strom 1 transportiert Daten von Station 1 zu Station 2, Strom 2 transportiert Daten von Station 2 zu Station 3, und Strom 3 transportiert Daten von Station 3 zu Station 1. Die an den jeweiligen Links angegebenen Kapazitäten dürfen da-

bei nicht überschritten werden, ansonsten seien alle Sender saturiert (d.h. bereit, mit beliebig hoher Datenrate zu senden). Die Menge aller gültigen Aufteilungen ist dann gegeben durch

$$X = \{(x_1, x_2, x_3) \in \mathbb{R}^3 \mid x_1 \in [0; 4 - x_3] \wedge x_2 \in [0; 2 - x_3] \wedge x_3 \in [0; 2]\}.$$

Dies lässt sich wie folgt begründen: Weil der Bottleneck-Link für Strom 3 die Kapazität 2 hat, muss zunächst $x_3 \in [0; 2]$ sein. Da auch Strom 2 diesen Bottleneck-Link passieren muss, folgt daraus, dass die Bandbreite, die für Strom 2 übrig bleibt, durch $2 - x_3$ nach oben begrenzt ist. Da die Bandbreite von Strom 2 nach oben durch 2 begrenzt ist, bleibt auf Link (4, 2) eine Restbandbreite von mindestens 4 für Strom 1 übrig. Deshalb ist Link (1, 4) der Bottleneck-Link für Strom 1, der mit Strom 3 geteilt wird, so dass für Strom 1 eine Restbandbreite von $4 - x_3$ verbleibt.

Bei der Aufteilung beliebiger Ressourcen ist es wünschenswert, dies nicht nur in irgendeiner gültigen, sondern in einer „fairen“ Art und Weise zu tun. Insbesondere ist zu beachten, dass die Aufteilung mit der höchsten Gesamtmenge an Ressourcen häufig besonders unfair ist, weil solche Instanzen, die in gewisser Weise ungünstige Voraussetzungen haben, gar nicht bedacht werden. Dies gilt auch für das angeführte Beispiel, in dem die Aufteilung (4, 2, 0) in der Summe zwar die meisten Ressourcen vergibt, Strom 3 aber unberücksichtigt lässt.

Die einfachste Form der Fairness ist die *strikte Fairness*, die allen Instanzen die gleiche, größtmögliche Menge an Ressourcen zuordnet. Für das in Abbildung 2.13 dargestellte Beispiel ist dies die Aufteilung (1, 1, 1), die sich daraus ergibt, dass sich die Ströme 2 und 3 den Link mit der Kapazität 2 teilen, der in diesem Beispiel der stärkste Engpass ist. Der offensichtliche Nachteil der strikten Fairness ist die Tatsache, dass Einschränkungen, die nur einen Teil aller Instanzen betreffen, auf alle Instanzen übertragen werden. Im beschriebenen Beispiel bleiben Ressourcen ungenutzt, die Strom 1 noch zugewiesen werden könnten.

Dies berücksichtigt das gängige Konzept der *Max-Min-Fairness* ([Jaf81]): Intuitiv wird eine Aufteilung von Ressourcen genau dann als max-min-fair bezeichnet, wenn einer Instanz ausgehend von dieser Aufteilung zusätzliche Ressourcen nur auf Kosten einer anderen Instanz zur Verfügung gestellt werden können, die gemäß der Ausgangsaufteilung auch schon nicht mehr Ressourcen als die erstgenannte zur Verfügung hatte. Formal sei $X \subseteq \mathbb{R}^n$ die Menge aller gültigen Aufteilungen von Ressourcen auf die n Instanzen. Eine Aufteilung $x = (x_1, \dots, x_n) \in X$ wird dann als max-min-fair bezeichnet, wenn für alle $y = (y_1, \dots, y_n) \in X$ gilt:

$$\exists i \in \{1, \dots, n\} : y_i > x_i \Rightarrow \exists j \in \{1, \dots, n\} : y_j < x_j \leq x_i.$$

Die max-min-faire Aufteilung im Beispiel in Abbildung 2.13 ist $x = (3, 1, 1)$. Wie oben beschrieben müssen sich die Ströme 2 und 3 die Ressourcen des stärksten Engpasses fair teilen, während Strom 1 mehr Ressourcen zugewiesen werden können, ohne dass dies die anderen beiden Ströme betrifft.

Daraus ergibt sich intuitiv eine konstruktive Vorgehensweise für die Bestimmung einer max-min-fairen Aufteilung, das sog. *Waterfilling* (wie es auch schon der in [Jaf81], Abschnitt V vorgestellte Algorithmus anwendet): Die vorhandenen Ressourcen werden so lange zu gleichen Teilen auf alle Instanzen verteilt, bis ein Engpass ausgeschöpft ist, so dass bestimmten

Instanzen keine weiteren Ressourcen mehr zugewiesen werden können. Diese Ressourcen werden dann blockiert, und das Waterfilling wird für die verbleibenden Instanzen fortgesetzt, bis der nächste Engpass ausgeschöpft ist. Dies wird so lange wiederholt, bis alle Instanzen blockiert sind.

Eine andere Fairness-Definition ist die *Utility-Fairness* [KMT98]. Dieser Fairness-Begriff beruht auf der Annahme, dass ein konkreter funktionaler Zusammenhang zwischen der Menge an zugeteilten Ressourcen und dem für die Anwendung resultierenden Nutzen besteht. Eine Aufteilung wird nach dieser Definition genau dann als fair bezeichnet, wenn sie den Gesamtnutzen maximiert.

Der funktionale Zusammenhang kann für die verschiedenen Instanzen unterschiedlich sein, und mehrere Beispiele werden in [She95] aufgeführt. Dieser Fairness-Begriff ist damit in erster Linie für Szenarien geeignet, in denen konkrete Annahmen über die Art der Anwendungen gemacht werden können. Dies ist in der vorliegenden Arbeit jedoch nicht der Fall, so dass sich hier die Max-Min-Fairness als Grundlage von Ressourcenaufteilungen anbietet.

3 Annahmen, Ziele und Anforderungen

3.1 Beschreibung der Anwendungsszenarien

Wie in Kapitel 1 ausgeführt, ist zu erwarten, dass in Zukunft immer mehr tragbare Geräte wie Mobiletelefone oder PDAs mit einer gewissen Rechenleistung und einer WLAN-Schnittstelle anzutreffen sein werden. Ist eine genügend hohe Anzahl solcher Geräte in ausreichender Dichte vorhanden, beispielsweise in Fußgängerzonen oder auf Massenveranstaltungen, können sie spontan ein Netz bilden, das selbstkonfigurierend und unabhängig von jeglicher Infrastruktur ist und das die Teilnehmer für Anwendungen wie Chat-Foren, File-Sharing oder Online-Spiele nutzen können. In der vorliegenden Arbeit wird untersucht, wie in solchen Netzen durch die Regelung der Sendeleistungen eine Steigerung des Spatial-Reuse und damit der Netzkapazität bewirkt werden kann.

Dabei wird vorausgesetzt, dass zumindest einige der beteiligten Geräte die Fähigkeit zur Regelung ihrer Sendeleistung haben, da sie technisch nicht schwierig umzusetzen und daher kaum mit Mehrkosten verbunden ist. Es ist aber unrealistisch, darüber hinaus besondere Anforderungen an die verwendete Hardware zu stellen, da zu erwarten ist, dass sich in einem solch heterogenen Umfeld durch die Konkurrenz verschiedener Hersteller untereinander vor allem preisgünstige Geräte durchsetzen werden. Außerdem wird angenommen, dass die Geräte einen gemeinsamen Übertragungskanal verwenden und dass deren Antennen omnidirektional sind. Wie im vorangegangenen Kapitel beschrieben, ist es recht kompliziert, ein drahtloses Netz, in dem die Geräte mobil sein können und nicht zwingend direkten Funkkontakt zueinander haben müssen, mit einer Multi-Channel-Technologie oder mit gerichteten Antennen zu realisieren, so dass davon auszugehen ist, dass die in der vorliegenden Arbeit betrachteten Szenarien häufig anzutreffen sein werden.

3.2 Zusammenhang zwischen Sendeleistungsregelung und Topologieformung

Die Art der Sendeleistungsregelung, wie sie in dieser Arbeit untersucht wird, ordnet während einer bestimmten Zeitspanne jedem Gerät eine eindeutige Sendeleistung zu, die es für alle Übertragungen unabhängig von deren Empfängern verwendet. Es werden zunächst also

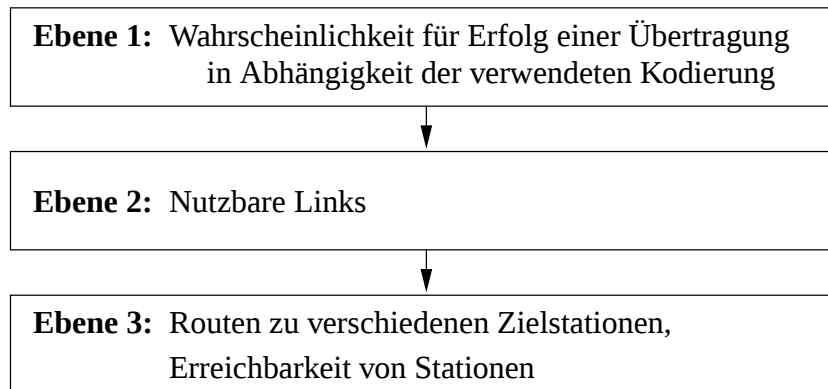


Abbildung 3.1: Einfluss der Sendeleistungsregelung auf verschiedenen Ebenen des OSI-Referenzmodells

Stellgrößen verändert, die eindeutig auf der Bitübertragungsebene (dem Physical-Layer) des OSI-Referenzmodells [Int94] anzusiedeln sind. Damit ändern sich auch Eigenschaften der Datenübertragung auf Bitübertragungsebene, nämlich die Wahrscheinlichkeit dafür, dass ein Empfänger ein auf bestimmte Weise kodiertes Funksignal eines Senders erfolgreich dekodieren kann.

Diese Wahrscheinlichkeiten wirken sich wiederum auf höhere Protokollschichten aus (siehe Abbildung 3.1). Auf der Verbindungsebene (dem Data-Link-Layer) wird dadurch beeinflusst, welche Links nutzbar sind. Durch die Festlegung der Sendeleistungen wird also die Topologie des Netzes geformt und damit das Routing beeinflusst, das auf Vermittlungsebene (dem Network-Layer) arbeitet.

Sendeleistungsregelung und Topologieformung müssen allerdings nicht so eng miteinander verknüpft sein. In vielen Arbeiten (z.B. [Bam98], [MBH01]) wird unter Sendeleistungsregelung verstanden, dass jede Übertragung gezielt mit einer günstigen Sendeleistung durchzuführen ist, die einerseits hoch genug ist, damit die Signalstärke am Empfänger zur Dekodierung ausreicht, die andererseits aber möglichst gering ist, damit ein hoher Spatial-Reuse erzielt werden kann. Diese Form der Sendeleistungsregelung ist nicht Gegenstand der vorliegenden Arbeit, da sie als zusätzliche, von der Topologieformung unabhängige Optimierung gesehen werden und mit ihr koexistieren kann: Wenn die Topologie durch das Festlegen einer maximalen Sendeleistung für jedes Gerät festgelegt ist, können einzelne Übertragungen mit geringerer Leistung durchgeführt werden, um dadurch eine zusätzliche Steigerung der Netzkapazität zu bewirken.

Die Topologieformung an sich ist allerdings unabdingbar, da es für einen sparsamen Umgang mit der verfügbaren Bandbreite sinnvoll ist, Links nicht zu verwenden, die eine hohe Sendeleistung erfordern. Der Broadcast-Verkehr, der in drahtlosen Netzen vor allem zur Signalisierung eine wichtige Rolle spielt, verbraucht dann weniger (oder zumindest nicht mehr) Ressourcen, da die gleichen Broadcasts durchgeführt werden, diese jedoch weniger sendewillige Geräte blockieren. Auch Unicast-Datenströme können unter gewissen Voraussetzungen

einen höheren Durchsatz erzielen, wenn für einen höheren Spatial-Reuse längere Pfade und damit mehr Übertragungen in Kauf genommen werden, wie auch in den folgenden Kapiteln 6 und 8 bestätigt wird.

3.3 Anforderungen an Verfahren zur Sendeleistungsregelung

Ein Mechanismus zur Sendeleistungsregelung, der in einem Umfeld eingesetzt werden soll, wie es in Abschnitt 3.1 beschrieben ist, muss bestimmte Anforderungen erfüllen, die im Folgenden erläutert werden.

Unabhängigkeit von Hardware und Umgebung Lösungen, die auf speziellen Fähigkeiten verwendeter Hardware oder auf spezifischen Eigenschaften der Einsatzumgebung aufbauen, sind nur für besondere Einsatzzwecke interessant, für die Geräte eigens hergestellt werden. Die Szenarien, die im Rahmen dieser Arbeit betrachtet werden, sind jedoch bewusst sehr unspezifisch als Anhäufung von Geräten definiert, die bestimmte Kommunikationsprotokolle beherrschen und damit spontan ein Netz bilden können. Wie in Abschnitt 3.1 bereits erläutert, können in einem solchen Umfeld keine besonderen Fähigkeiten von der eingesetzten Hardware erwartet werden. Außerdem ist nicht davon auszugehen, dass alle Geräte im Netz dieselbe Funktionalität haben; insbesondere können die am Netz beteiligten Geräte unterschiedliche Mengen von Sendeleistungen unterstützen.

Ebenso ist es selbstverständlich, dass eine allgemein akzeptierte Lösung in einer möglichst großen Vielfalt an Szenarien einsetzbar sein muss. Es muss damit gerechnet werden, dass das Netz sich in einer geschlosseneren Umgebung befindet, also innerhalb von Gebäuden oder in dicht bebauten Gebieten. In solchen Umgebungen ist die Wahrscheinlichkeit dafür, dass Radiowellen entlang der Sichtlinie zwischen Sender und Empfänger übertragen werden können, im Vergleich zu offeneren Umgebungen relativ gering. Außerdem beinhaltet diese Forderung auch, dass die Netze dynamisch sein können in dem Sinn, dass die beteiligten Geräte bewegt und spontan ein- oder ausgeschaltet werden.

Es kann allerdings davon ausgegangen werden, dass in den meisten Fällen keine allzu hohe Dynamik vorliegt, da die Bewegungsgeschwindigkeit der Geräte in den eingangs skizzierten Szenarien höchstens der von Fußgängern entspricht. Außerdem werden viele der Geräte auch statisch sein, da Menschen dazu neigen, bei der Interaktion mit Anwendungen auf ihrem Endgerät nicht in Bewegung zu sein, sondern stehenzubleiben oder es sich im Sitzen bequem zu machen.

Güte der erzeugten Topologien Das Ziel der Sendeleistungsregelung ist die Steigerung der Netzkapazität. Neben Eigenschaften, die eine unmittelbare Auswirkung auf die Netzkapazität haben, sind aber noch weitere Bedingungen zu beachten.

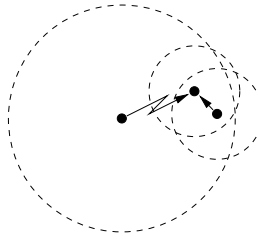


Abbildung 3.2: Verschärfung des Hidden-Node-Problems durch stark verschiedene Sendeleistungen in räumlicher Nähe zueinander

Obwohl es sinnvoll ist, die Sendeleistungen an lokale Gegebenheiten des Netzes anzupassen, muss gefordert werden, dass relativ nahe zueinander positionierten Geräte möglichst nicht zu unterschiedliche Sendeleistungen zugewiesen werden. Dadurch können Mechanismen zum Carrier-Sensing im Rahmen ihrer Möglichkeiten effektiv bleiben und das Hidden-Node-Problem wird nicht unnötig verstärkt. Ansonsten passiert es leicht, dass ein Gerät, das mit relativ hoher Leistung sendet, weder mit Hilfe des physikalischen noch mit Hilfe des virtuellen Carrier-Sensing eine laufende Übertragung wahrnehmen kann, die es durch eine Übertragung seinerseits stören würde (dies wird in Abbildung 3.2 veranschaulicht).

Aus diesem Grund ist auch zu vermeiden, dass Geräte in regelmäßigen Zeitabständen Kontrollnachrichten mit maximaler Leistung versenden, um sich ein Bild von möglichen Verbindungen zu machen, die bei den aktuell eingestellten Sendeleistungen nicht bestehen. Dies wäre notwendig, um eine Reihe von Verfahren für statische Netze an den Betrieb in dynamischen Netzen anzupassen, wie in Abschnitt 4.2 detaillierter ausgeführt wird.

Minimierung des Protokoll-Overheads Das vorrangige Ziel ist es, den Netzteilnehmern eine hohe Übertragungskapazität zu bieten. Deshalb muss darauf geachtet werden, dass der Bandbreitengewinn nicht durch den Overhead kompensiert (oder gar überkompensiert) wird, den das Protokoll zur verteilten Sendeleistungsregelung erzeugt. Damit scheiden Ansätze aus, die eine globale Sicht des Netzes erfordern, denn gerade in großen Netzen wäre ein regelmäßiges Fluten von Topologie-Informationen viel zu aufwändig.

Daraus ergibt sich, dass die Sendeleistungsregelung dezentral erfolgen muss. Dabei ist zu bedenken, dass dazu entweder sichergestellt sein muss, dass alle Teilnehmer die gleiche Konfiguration verwenden, oder dass die Funktionalität des Algorithmus nicht wesentlich dadurch beeinträchtigt wird, dass Geräte in einem Netz auch unterschiedlich konfiguriert sein können. Ersteres ist allerdings höchstens in Netzen mit geschlossener Benutzergruppe möglich und selbst dort schon organisatorisch schwierig, wenn es sich um mehr als nur einige wenige Geräte handelt. Für eine allgemeinere Einsetzbarkeit ist es daher zwingend erforderlich, dass auch bei unterschiedlicher Konfiguration der Protokollinstanzen die Funktionalität des Algorithmus erhalten bleibt.

Die in dieser Arbeit betrachteten Netze haben schließlich von sich aus eine gewisse Dynamik, die aus den Bewegungen und dem spontanen An- und Abschalten der Geräte resultiert,

und für drahtlose Multihop-Netze entwickelte Routingprotokolle sind auf den Einsatz unter solchen Umständen ausgelegt. Dennoch ist zu beachten, dass diese Dynamik durch die Regelung der Sendeleistungen und die daraus resultierenden Veränderungen der Netztopologie steigt. Das Routingprotokoll sollte nicht zu sehr durch unnötig viele dieser Änderungen beansprucht werden. Insbesondere das Verringern von Sendeleistungen kann zum Wegfallen von Links führen, die gerade zur Datenübertragung verwendet werden, und damit zusätzlichen Kontrollverkehr zum Aufbau neuer Routen und ggf. Datenverluste verursachen.

4 Verwandte Arbeiten

4.1 Auswirkungen der Sendeleistungsregelung

4.1.1 Analytische Arbeiten

Schon im Rahmen der Forschung zu PRNs, die mittlerweile drei Dekaden zurückreicht, finden sich Arbeiten zur Auswirkung der Sendeleistungsregelung auf die Kapazität drahtloser Multihop-Netze. Ein guter Überblick über die Forschung aus dieser Zeit ist [KS87].

Besonders erwähnenswert sind in diesem Zusammenhang die Arbeiten [KS78] und [TK84], die für verschiedene Arten des Medienzugangs herleiten, welche durchschnittliche Nachbarnzahl die Stationen in Topologien mit maximaler Transportkapazität haben. Dabei wird angenommen, dass alle Stationen gleichverteilt positioniert sind und die gleiche Sendereichweite haben. Weiterhin wird ein unendlich großes Netz angenommen, über das der Datenverkehr homogen verteilt ist. Je nach Art des Medienzugangs und Netzmodell erweisen sich Nachbarnzahlen im Bereich von 6 bis 8 als optimal. Diesen Arbeiten liegt die Erkenntnis zugrunde, dass die Größe der Fläche, die durch eine Übertragung abgedeckt und damit blockiert wird, quadratisch von der Sendereichweite abhängt, die Pfadlänge zwischen einem Sender und einem Empfänger hingegen nur reziprok linear. Dadurch begründet sich die positive Auswirkung des gesteigerten Spatial-Reuse auf die Transportkapazität, der die erhöhte Anzahl von Übertragungen überkompensiert, die mit verringerten Sendeleistungen einhergeht.

In [GK00] wird die Transportkapazität von Netzen mit einheitlichen Sendereichweiten analysiert. Dazu werden Netze aus n zufällig positionierten Stationen untersucht, in denen jede Station an eine zufällig gewählte Zielstation sendet. Der Durchsatz, den jeder einzelne Strom so erreichen kann, hängt davon ab, wie die Stationen positioniert sind und wie die Quelle-Senke-Paare gewählt sind, so dass [GK00] dazu keine konkreten Aussagen macht. Es wird aber bewiesen, dass die Wahrscheinlichkeit dafür, dass jeder Strom einen Durchsatz von $c \cdot W / \sqrt{n \log n}$ erreichen kann, für $n \rightarrow \infty$ gegen 1 konvergiert, wobei W die Linkbandbreite angibt und $c > 0$ konstant ist. Je größer die Stationszahl eines Netzes also wird, umso wahrscheinlicher ist es, dass der Anteil der Linkbandbreite, die eine Station für den selbst generierten Verkehr nutzen kann, in $O(1/\sqrt{n \log n})$ ist. Dies wird sowohl für das Protokollmodell als auch für das physikalische Modell gezeigt (siehe Abschnitt 2.4.1).

In [GC04] wird der Kommentar der Autoren in [GK00] aufgegriffen, dass Sendereichweiten so klein wie möglich sein sollten, da der Spatial-Reuse stärker ins Gewicht fällt als die Länge der Routen. Deshalb wird der Versuch unternommen, die Kapazität für den Fall zu quantifizie-

ren, dass jeder Station eine günstige individuelle Sendeleistung zugewiesen wird. Hier wird der vermeintlich günstigste Fall betrachtet, dass die Netztopologie ein minimaler Spannbaum ist. Die Autoren kommen dann zu dem erstaunlichen Ergebnis, dass die Bandbreite, die einer Station zu Verfügung steht, mit steigender Stationszahl konstant ist.

Bei genauerer Betrachtung ist dieses Ergebnis jedoch wenig aussagekräftig, da es auf der Abschätzung einer oberen Schranke beruht. Während in [GK00] auch eine untere Schranke der gleichen Größenordnung bewiesen wird, fehlt eine solche Ergänzung in [GC04]. Insbesondere verletzen auf minimalen Spannbäumen basierende Topologien eine Bedingung, die im konstruktiven Beweis der unteren Kapazitätsschranke in [GK00] gefordert wird. Informell ist dies die Bedingung, dass die Sendereichweiten der Stationen hinreichend hoch sind, so dass der Dehnungsfaktor (siehe Abschnitt 2.4.2) beschränkt ist, die Länge eines Pfads die Distanz zwischen Sender und Empfänger also um nicht mehr als einen konstanten Faktor übersteigt. Dies trifft für minimale Spannbäume jedoch nicht zu; in Abschnitt 6.5 wird gezeigt, dass der Dehnungsfaktor gerade in Topologien, die auf minimalen Spannbäumen beruhen, extrem hohe Werte annimmt, so dass diese Topologien nur eine geringe Transportkapazität aufweisen.

Weiterhin versucht [GC04], den Kompromiss zwischen zwei gegenläufigen Folgen verringerter Sendeleistungen genauer zu beleuchten: Während die Netzkapazität an sich erhöht wird, erzeugt das Routingprotokoll eine größere Menge an Kontrollverkehr, da die Linklebenszeiten angesichts der Mobilität der Endgeräte kürzer werden (die Geräte sind kürzer in gegenseitiger Sendereichweite). Dieser Zusammenhang wird auch in den Auswertungen in Kapitel 8 untersucht.

In [HM04a] wird die Kapazität von drahtlosen Multihop-Netzen ausgehend vom erwarteten SINR in vorgegebenen Topologien abgeschätzt, die sich durch eine Positionierung der Stationen an den Schnittpunkten eines Gitters aus gleichseitigen Dreiecken ergibt. Aus dem SINR lässt sich mit Hilfe der Shannon-Formel für die Kanalkapazität ([Sha49], [Rap96]) die maximale entlang eines Links erzielbare Datenrate berechnen (in Abhängigkeit von der Signalbandbreite). In den Auswertungen wird hauptsächlich der Einfluss der Stationszahl, nicht aber der Sendeleistungen untersucht. Es wird lediglich die Vermutung geäußert, dass der Spatial-Reuse in vielen Fällen stärker ins Gewicht fallen kann als die Pfadlängen zwischen Sendern und Empfängern, was auch anhand eines Beispiels illustriert wird. Eine systematische Untersuchung der Wechselwirkung zwischen diesen beiden Größen wird aber nicht vorgenommen.

Die vorgestellten Arbeiten bieten zum Teil interessante Einblicke in die Auswirkungen der Sendeleistungsregelung auf die Kapazität des resultierenden Netzes. In den meisten Fällen sind die Ergebnisse aber nur für Netze gültig, in denen alle Stationen mit der gleichen Leistung senden. Ein Vergleich von unterschiedlichen Strategien zur Sendeleistungsregelung auf analytischer Ebene scheint dagegen kaum realisierbar zu sein.

In [HP05] werden Argumente zusammengefasst, die gegen die Verringerung von Sendeleistungen zur Steigerung der Netzkapazität sprechen. Die Existenz dieser begründeten Argumente bezeugt, dass der Zusammenhang zwischen den Sendeleistungen der Geräte und der Netzkapazität komplex ist und dass eine Verringerung der Sendeleistungen nicht immer zu einer Kapazitätssteigerung führt. Die in [HP05] vorgebrachten Argumente beruhen aber auf be-

stimmten Annahmen, so dass die Autoren als Fazit auch betonen, dass Kommunikationspfade über Links, die möglichst große Distanzen abdecken, „in vielen Netzen“ nicht schlechter sind als Kommunikationspfade mit mehr Links, die kürzere Distanzen abdecken. Die vorliegende Arbeit kommt zu dem gleichen Schluss und identifiziert Eigenschaften, die Netze, in denen eine Verringerung der Sendeleistungen sich positiv auf die Kapazität auswirkt, von anderen Netzen unterscheiden, in denen dies nicht der Fall ist.

4.1.2 Eigenschaften konkreter Topologien

Wie in Abschnitt 2.4.4 erläutert, werden in der vorliegenden Arbeit zwei Arten der Netzkapazität unterschieden. In Abschnitt 4.1.2.1 werden Arbeiten aus dem Bereich der Broadcast-Kapazität vorgestellt und im darauf folgenden Abschnitt 4.1.2.2 Arbeiten aus dem Bereich der Transportkapazität.

Ein völlig anderer Bewertungsansatz für Topologien, der keiner der beiden Kategorien angehört, wird in [DRL03] beschrieben. Die (gerichteten) Kanten einer Topologie werden mit Gewichten versehen, die die Wahrscheinlichkeit einer erfolgreichen Übertragung entlang der Kante innerhalb eines Zeitslots unter bestimmten Annahmen über die Konkurrenz beim Medienzugang angeben. Weiterhin wird für jeden Knoten v eine Kante (v, v) hinzugefügt, die mit der Wahrscheinlichkeit gewichtet ist, dass diese Station innerhalb eines Zeitslots keine Übertragung durchführt. Auf diese Weise entsteht der Übergangsgraph einer Markov-Kette, die hinsichtlich ihrer Konvergenzgeschwindigkeit auf die eindeutige stationäre Wahrscheinlichkeitsverteilung untersucht werden kann. Die Autoren von [DRL03] argumentieren, dass diese Konvergenzgeschwindigkeit ein Maß für die „Effizienz“ einer Topologie sei. Es soll ein Zusammenhang zur Zeitspanne bestehen, die benötigt wird, bis ein proaktives Routingprotokoll bei allen Stationen konsistente Routingtabellen erzeugt hat, es wird aber auch ein Bezug zur Netzkapazität hergestellt.

Genaue Untersuchungen in [Pes04] zeigen jedoch, dass ein solcher Bezug nicht gegeben ist, sondern dass die besagte Konvergenzgeschwindigkeit mit steigenden Sendereichweiten grundsätzlich ansteigt. Dies ist auch dann der Fall, wenn andere, realitätsnähere Modellierungen für Konkurrenz beim Medienzugang verwendet werden als in [DRL03].

4.1.2.1 Broadcast-Kapazität

Wie in Abschnitt 2.4.4 erläutert wurde, bezeichnet die *Broadcast-Kapazität* in der vorliegenden Arbeit den Durchsatz, der erzielbar ist, wenn Datenpakete von allen Stationen im Netz genau einmal weitergeleitet werden (siehe das in Abschnitt 2.3.1 beschriebene *Fluten*). Im Vergleich zur Transportkapazität, auf die im folgenden Abschnitt genauer eingegangen wird, ist die Broadcast-Kapazität vergleichsweise einfach zu quantifizieren: Es muss lediglich ein Übertragungsschedule (siehe Abschnitt 2.4.3) gefunden werden, der in möglichst kurzer Zeit jeder Station genau eine Übertragung erlaubt, die von allen Nachbarn eines Senders empfangen werden kann (d.h. dies wird nicht durch Interferenz anderer Stationen verhindert).

Dieses Problem wird in vielen Fällen dadurch gelöst, dass ein *Distance-2-Vertex-Colouring* vorgenommen wird. Eine solche Färbung hat die Eigenschaft, dass ein potenzieller Empfänger einer Übertragung im Topologiegraphen nicht mit einem weiteren Knoten verbunden ist, der zur gleichen Zeit sendet. Diese Vorgehensweise ist jedoch nur dann sinnvoll, wenn die Topologie keine unidirektionalen Links enthält, wenn also einheitliche Sendereichweiten modelliert werden, wie in Abschnitt 2.4.3 erläutert wird.

In [ZK03] wird versucht, die Dauer (engl. *settling time*) einer einzelnen Flutung abzuschätzen. Für die Optimierung dieses Werts ist unter anderem jedoch auch die Paketverzögerung relevant, während diese für die Broadcast-Kapazität keine Rolle spielt. Dies äußert sich darin, dass in [ZK03] beobachtet wird, dass eine Erhöhung der einheitlichen Sendereichweite über den Wert hinaus, der die Konnektivität des Netzes gewährleistet, die Dauer einer Flutung reduziert, während die Broadcast-Kapazität im oben beschriebenen Sinn monoton mit steigenden Sendereichweiten fällt. Insofern ist der Schluss der Autoren, dass der durch Flutung erzielbare Durchsatz umgekehrt proportional zur Dauer einer Flutung ist, etwas zu voreilig, da unberücksichtigt bleibt, dass mehrere Flutungsvorgänge zeitlich überlappend stattfinden können.

4.1.2.2 Transportkapazität

Die *Transportkapazität* wird in Abschnitt 2.4.4 als der maximale erzielbare, aggregierte Ende-zu-Ende-Durchsatz von Unicast-Datenströmen definiert. Diese Größe hängt daher auch maßgeblich davon ab, welche konkreten Quelle-Senke-Paare dazu gewählt werden.

Sollen verschiedene konkrete Topologien hinsichtlich ihrer Transportkapazität verglichen werden, so muss der Vergleich auf denselben Verkehrsmustern basieren. Es ist daher nicht sinnvoll, die Transportkapazität einer Topologie als den maximal erzielbaren Ende-zu-Ende-Durchsatz über alle möglichen Verkehrsmuster zu definieren. In diesem Fall würde für jede Topologie ein spezielles Verkehrsmuster zur Geltung kommen und damit zu einer optimistischen Bewertung führen.

In [JS03] wird ein Verfahren zur Errechnung der Kapazität konkreter Topologien vorgestellt, das auf Annahmen beruht, die die Anwendungsgebiete des Verfahrens stark eingrenzen. Zunächst wird angenommen, dass alle Datenströme eine gemeinsame Senke haben, was beispielsweise dann der Fall ist, wenn alle Benutzer ausschließlich auf einen Server oder Zugangspunkt zu anderen Netzen zugreifen. Weiterhin erläutern die Autoren die Schwierigkeit, Fairness in geeigneter Weise durchzusetzen und kommen zu dem Schluss, dass strikte Fairness eingehalten werden muss, dass also allen Datenströmen die gleiche Bandbreite zugewiesen wird. Aufgrund dieser Fairness-Annahme lässt sich in einer vorgegebenen Topologie leicht der stärkste Engpass identifizieren, der wiederum die verfügbare Bandbreite für alle Ströme bestimmt.

Falls alle Datenströme die gleiche Senke haben, besteht in der Regel kein Unterschied zwischen strikter Fairness und Max-Min-Fairness, da die Links, die zur Senke führen, mit hoher Wahrscheinlichkeit auch zum stärksten Engpass gehören, der somit alle Ströme betrifft.

Werden jedoch beliebige Quelle-Senke-Paare zugelassen, so lässt sich unter Einhaltung der Max-Min-Fairness ein deutlich höherer Gesamtdurchsatz erzielen, indem Ressourcen, die unter Einhaltung strikter Fairness nicht genutzt werden, auf die Datenströme verteilt werden, die diese nutzen können. In diesem Fall bestimmen die Datenraten der einzelnen Ströme jedoch, wo sich Engpässe befinden, die unter der Annahme von Max-Min-Fairness wiederum die Datenraten der Ströme beeinflussen. Wegen dieser wechselseitigen Abhängigkeit ist der in [JS03] beschriebene Ansatz daher nicht mehr anwendbar.

Ein Verfahren, das die Berechnung der Transportkapazität von Topologien für beliebige Quelle-Senke-Paare ermöglicht, findet sich in [JPPQ03]. Indem die Datenübertragung als Flussproblem aufgefasst wird, kann der Gesamtdurchsatz als maximaler Fluss berechnet werden, wobei zusätzliche Einschränkungen beachtet werden müssen, die die drahtlose Hardware und ggf. Fairness-Bedingungen modellieren. Obwohl der maximale Fluss aus Effizienzgründen nur approximiert wird, ist die Laufzeit des vorgestellten Verfahrens selbst für kleinere Netze sehr hoch, weshalb die Autoren auch lediglich Ergebnisse für Netze mit weniger als 50 Stationen präsentieren.

Eine ähnliche Vorgehensweise wird in [KN03] verfolgt. Um den maximalen Durchsatz eines Datenstroms zu ermitteln, wird zunächst ein maximaler Fluss approximiert, der allerdings eine Obergrenze für die erzielbare Datenrate darstellt, da die Einschränkungen hier notwendig und nicht hinreichend für die Realisierbarkeit sind. Um den tatsächlich realisierbaren Durchsatz zu bestimmen, wird eine Kantenfärbung des Multigraphen durchgeführt, in dem jede Kante eine Übertragung im zuvor bestimmten maximalen Fluss repräsentiert. Für mehrere Datenströme wird ein ähnliches Verfahren präsentiert, um zu überprüfen, ob ein vorgegebener Vektor von Datenraten erzielbar ist. Weil in diesem Fall die Datenraten konkret vorgegeben werden müssen, wären viele Iterationen nötig, um den maximalen Durchsatz mehrerer Ströme unter Einhaltung gewisser Fairness-Bedingungen zu berechnen.

Außerdem werden in [KN03] explizit Multi-Channel-Netze modelliert: Die einzige Einschränkung ist die, dass Stationen zu einem bestimmten Zeitpunkt nur an einer Übertragung beteiligt sein können. Die Erweiterbarkeit auf andere Modelle wird zwar angesprochen, jedoch würde das beschriebene Verfahren dadurch noch rechenaufwändiger. Die präsentierten Ergebnisse beziehen sich aber wieder auf Netze mit sehr geringer Stationszahl, nämlich höchstens 30 Stationen.

4.1.3 Zusammenfassung und Motivation neuer Verfahren zur Kapazitätsabschätzung

Es wurde schon vergleichsweise früh damit begonnen, die Leistungsfähigkeit bestimmter Topologieklassen auf analytischer Ebene zu untersuchen. Diese Arbeiten sind für ein grundlegendes Verständnis der Leistungsfähigkeit drahtloser Multihop-Netze wichtig. Die Ergebnisse können auch dazu dienen, für Szenarien, deren Eigenschaften im Voraus bekannt sind, günstige Parameter zu wählen. Sie bieten aber keine Vergleichsmöglichkeit für Topologien, die aus unterschiedlichen Strategien zur Topologieformung resultieren.

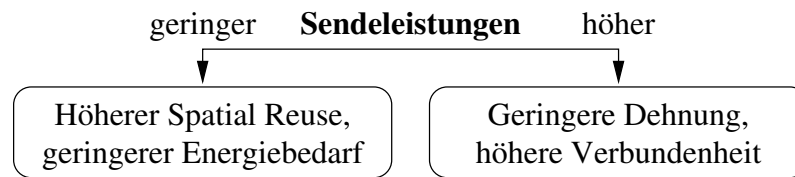


Abbildung 4.1: Sendeleistungsregelung als Abwägung zwischen entgegengesetzten Zielen

Zur Charakterisierung der Broadcast-Kapazität wie auch deren Quantifizierung in gegebenen Topologien gibt es nur wenige Veröffentlichungen. Der grundsätzliche Zusammenhang zwischen den verwendeten Sendeleistungen und der resultierenden Broadcast-Kapazität ist nämlich relativ klar: Je größer die verwendeten Sendeleistungen sind, desto geringer ist der Spatial-Reuse und damit auch die Broadcast-Kapazität. Das Thema stellt somit keine besondere Herausforderung dar (im Gegensatz zur Messung der Transportkapazität). Eine einfache und allgemeingültige Lösung, die in der vorliegenden Arbeit verwendet wird, wird in Abschnitt 5.2 präsentiert.

Zur Abschätzung der Transportkapazität konkreter Topologien sind bislang nur wenige Verfahren vorgestellt worden, und diese sind nur bedingt praktisch anwendbar, weil sie entweder lediglich auf bestimmte Verkehrsmuster ausgelegt oder aber zu rechenaufwändig sind, als dass sie mit vertretbarem Aufwand für große Netze durchführbar wären. Deshalb wurde im Rahmen der vorliegenden Arbeit ein eigenes Verfahren entwickelt, das auch für größere Topologien mit mehreren tausend Stationen eine Abschätzung des Durchsatzes von Datenströmen zwischen beliebigen Quelle-Senke-Paaren praktisch erlaubt ([Pes04], [PGdWM04]). Dieses Verfahren wird in Abschnitt 5.3 beschrieben.

4.2 Ansätze zur Topologieerzeugung

In den folgenden Abschnitten werden verschiedene Verfahren zur Erzeugung von als günstig erachteten Netztopologien vorgestellt. Der hier gegebene Überblick beschränkt sich auf Verfahren, die der in Abschnitt 3.2 beschriebenen Regelung der Sendeleistung der einzelnen Stationen in Single-Channel-Netzen dient. Die anderen Arbeiten sind im Kontext der vorliegenden Arbeit nicht relevant, weil sie auf Voraussetzungen aufbauen, die in den hier betrachteten Netzen nicht gegeben sind.

Die vorgestellten Arbeiten verfolgen zwei Ziele. Eine Erhöhung des Spatial-Reuse soll eine Steigerung der Netzkapazität bewirken, und eine Verringerung der Sendeleistungen soll den Energiebedarf der kleinen, batteriebetriebenen Geräte senken. Da die Steigerung des Spatial-Reuse ebenfalls die Verringerung der Sendeleistungen erfordert, wird häufig keine klare Abgrenzung zwischen diesen Zielen vorgenommen, zumal in den meisten Fällen nur die Intuition als Begründung dafür dient, dass ein erklärtes Ziel erreicht wird: Eine Senkung der Sendeleistungen unter die von der Hardware vorgegebenen Maximalwerte kann den mittleren Energiebedarf einer Übertragung und auch die mittleren Knotengrade der resultierenden

Netztopologien verringern. Letzteres bewirkt, dass die Konkurrenz um das gemeinsame Übertragungsmedium sinkt und der Spatial-Reuse damit steigt.

In beiden Fällen muss jedoch beachtet werden, dass eine geeignete Wahl von Sendeleistungen ein Kompromiss zwischen entgegengesetzten Zielen ist (siehe Abbildung 4.1): Die Verringerung der Sendeleistungen kann zunächst bewirken, dass die resultierende Netztopologie in mehrere Zusammenhangskomponenten zerfällt, so dass zwischen bestimmten Stationspaaren gar keine Verbindung mehr möglich ist. Daher liegt die Forderung nahe, dass die Konnektivität gewahrt sein soll, die das Netz mit maximalen Sendeleistungen bietet. Insbesondere ist es wünschenswert, dass die Topologie verbunden ist, sofern dies möglich ist.

Aber auch dann, wenn die Verbundenheit der Netztopologie erhalten bleibt, sind bei verringerten Sendeleistungen im Allgemeinen mehr Übertragungen erforderlich, um ein Paket von seiner Quelle zu seiner Senke zu transportieren. Insbesondere steigt in der Regel auch die Dehnung, d.h. der Quotient zwischen Pfadlänge und euklidischer Distanz (siehe Abschnitt 2.4.2). Deshalb finden in einigen Fällen explizit Graphtypen Beachtung, die Garantien auf die maximale Dehnung (den Dehnungsfaktor) geben können. Dadurch kann garantiert werden, dass nicht zu viel von dem Gewinn an Spatial-Reuse durch längere Ende-zu-Ende-Pfade wieder verloren geht.

Eine Reihe von Verfahren sind für statische Netze konzipiert und somit in erster Linie für Sensornetze interessant: Die Topologie wird in einer Initialisierungsphase festgelegt und ändert sich danach nicht. Um solche Algorithmen in dynamischen Netzen anwenden zu können, müssten in regelmäßigen Zeitabständen Entscheidungen getroffen werden, die eine Kenntnis aller möglichen Links erfordern, die mit den aktuell verwendeten Sendeleistungen eventuell nicht vorhanden sind. Dies würde wiederum erfordern, dass die Stationen in regelmäßigen Zeitabständen Kontrollpakete mit maximaler Sendeleistung versenden. Dies hätte vor allem in Netzen, in denen eigentlich sehr viel geringere Sendeleistungen ausreichen, sehr ungünstige Auswirkungen auf den laufenden Betrieb: Zum einen ginge viel Kapazität verloren, weil eine vergleichsweise hohe Anzahl von Stationen durch diese Kontrollpakete blockiert würde, und zum anderen könnte zumindest nur sehr schwierig sichergestellt werden, dass laufende Übertragungen ungestört bleiben, die der Sender des Kontrollpakets nicht wahrnehmen kann (siehe Abbildung 3.2).

Die meisten Ansätze lassen sich in zwei Kategorien einteilen. In der überwiegenden Zahl der Arbeiten wird ausschließlich anhand eines Distanzmaßes festgelegt, welche Stationen in der zu erzeugenden Topologie direkt verbunden sein sollten und welche nicht. Diese Gruppe von Ansätzen wird in Abschnitt 4.2.1 vorgestellt. Obwohl in vielen Fällen implizit davon ausgegangen wird, dass dieses Maß die euklidische Distanz zwischen den Positionen der Stationen in einer Ebene ist, können ebenso andere Maße verwendet werden, z.B. die beobachtete Linkqualität bei einer vorgegebenen Sendeleistung.

Weiterhin gibt es mehrere Ansätze, die als Eingabe die Positionen der Stationen in einem euklidischen Raum oder vergleichbare Information benötigen, da hier geometrische Eigenschaften wie z.B. der Winkel zwischen zwei Kanten eine Rolle spielen. Diese in der vorliegenden Arbeit als *geometrische Ansätze* bezeichnete Gruppe wird in Abschnitt 4.2.2 vorgestellt. Da-

mit verbundenen Nachteilen werden zu Beginn von Abschnitt 4.2.2 dargelegt.

Schließlich werden in Abschnitt 4.2.3 konkrete Verfahren vorgestellt, die sich keiner der beiden Kategorien zuordnen lassen.

4.2.1 Nur auf einem Distanzmaß basierende Ansätze

Die nun folgenden Ansätze beruhen auf der Definition eines *Nachbarschaftsgraphen* (engl. *proximity graph*), der eine Kante genau dann enthält, wenn die inzidenten Knoten ein bestimmtes Kriterium erfüllen, das mit Hilfe eines Distanzmaßes definiert ist. Die verschiedenen Definitionen von Nachbarschaftsgraphen stammen ursprünglich nicht aus der Forschung zu Rechnernetzen, sondern aus anderen Bereichen der Informatik, insbesondere der Forschung zu Clusteringmethoden.

Es ist zu bedenken, dass ein sinnvolles Distanzmaß enger mit der Dämpfung korrelieren sollte als die räumliche Entfernung, die oft nur ein schlechter Indikator für die Dämpfung eines Radiosignals ist, die bei einer Übertragung zu erwarten ist (siehe Abschnitt 2.1.3.1). Schließlich ist eine geringe Dämpfung entscheidend dafür, dass Links tatsächlich mit auch geringer Sendeleistung genutzt werden können.

Zur Vereinfachung wird im Folgenden angenommen, dass die Distanzen von einem beliebigen Knoten zu allen anderen Knoten eindeutig sind, dass also $d_{uv} \neq d_{uw}$ für alle paarweise verschiedenen $u, v, w \in S$. Kann dies nicht garantiert werden, bietet sich die Erweiterung des Distanzmaßes dahingehend an, dass bei gleicher Distanz beispielsweise die Knotenindizes eine Reihenfolge festlegen.

4.2.1.1 Max-Degree-Topologien

Mehrere Forschungsarbeiten präsentieren Algorithmen zur Sendeleistungsregelung, die auf der einfachen Grundidee beruhen, dass die Anzahl der Links jeder Station in einem bestimmten Bereich $[k_{\min}; k_{\max}]$ liegen soll. Diese Idee ist eine naheliegende Konsequenz aus den in Abschnitt 4.1.1 vorgestellten analytischen Arbeiten, die den Durchsatz in Topologien mit bestimmten durchschnittlichen Nachbarzahlen untersuchen und feststellen, dass zur Maximierung des Durchsatzes zumindest unter idealisierten Bedingungen eine optimale Nachbarzahl erreicht werden sollte.

Zur dezentralen Umsetzung dieser Grundidee bietet es sich an, den einfachen Zusammenhang auszunutzen, dass die Erhöhung der Sendeleistung im Allgemeinen zu einer höheren Anzahl bidirektionaler Links führt, während die Verringerung der Sendeleistung die Anzahl bidirektionaler Links reduziert. Eine Station erhöht ihre Sendeleistung also, wenn sie zu wenig Links hat, und verringert sie, wenn sie zu viele Links hat.

Dieses Verfahren hat zwei Vorteile. Zunächst ist das zugrunde liegende Distanzmaß implizit durch die Sendeleistung bestimmt, die notwendig ist, damit ein unidirektionaler Link von einer Station zu einer anderen besteht. Diese hängt wiederum von der Dämpfung ab, die ein

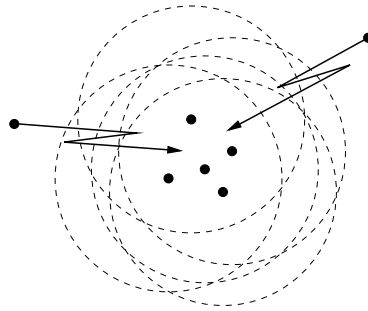


Abbildung 4.2: Schwäche des Max-Degree-Ansatzes ($k_{\max} = 4$)

Funktsignal erfährt. Es handelt sich also um ein gemäß den zu Beginn dieses Abschnitts 4.2.1 dargelegten Überlegungen praktisches Distanzmaß. Außerdem muss eine Station ihre Sendeleistung nicht ändern, so lange ihre Nachbarzahl im vorgegeben Zielbereich liegt. Es ist also nicht, wie für viele andere Kriterien zur Topologieformung erforderlich, dass eine Stationen Informationen darüber hat, welche anderen Links zu weiter entfernten Stationen möglich wären und welche Eigenschaften diese hätten.

Unter der Voraussetzung, dass die Sendeleistungen mit beliebig feiner Granularität eingestellt werden können und dass die Distanzen aller Nachbarn einer Station unterschiedlich sind, wird die Topologie eines statischen Netzes nach einer Einschwingphase feststehen. Die in dieser Arbeit verwendete Bezeichnung der *Max-Degree-Topologie* bezieht sich auf die Eigenschaft einer solchen Topologie, dass der maximale Knotengrad den Schwellwert $k_{\max}$ nicht überschreitet. Dies kann zu ungünstigen Situationen führen, wie Abbildung 4.2 illustriert: Hier liegen einige Stationen recht nahe beieinander, so dass sie mit relativ geringer Sendeleistung eine ausreichende Anzahl von Links aufbauen können. Sie werden ihre Sendeleistungen keinesfalls erhöhen, denn sie haben bereits genügend Links und würden sonst die maximale Nachbarzahl $k_{\max}$ überschreiten. Andere Stationen, die etwas weiter entfernt sind, müssen ihre Sendeleistungen unter Umständen immer weiter erhöhen, weil sie die Mindestlinkzahl $k_{\min}$ nicht erreichen können. Solche Stationen können vollständig vom Netz isoliert bleiben und senden dann regelmäßig mit maximaler Sendeleistung Beacons oder Route-Request-Nachrichten. Wenn mehrere Stationen in der Umgebung mit diesem Problem konfrontiert sind, können sie bidirektionale Links etablieren, die extrem hohe Sendeleistungen erfordern. In beiden Fällen ist es möglich, dass diese Stationen die laufenden Übertragungen anderer Stationen, die mit sehr viel geringeren Sendeleistungen durchgeführt werden, gar nicht wahrnehmen können, so dass sie diese durch ihre eigenen Übertragungen stören.

Der Max-Degree-Ansatz wurde zum ersten Mal in [RRH00] als *LINT-Algorithmus* (Local Information, No Topology) vorgestellt. Da von einer Hardware ausgegangen wird, die die Einstellung der Sendeleistung auf beliebige Werte innerhalb eines kontinuierlichen Bereichs erlaubt, wird basierend auf dem gewünschten Knotengrad, dem aktuellen Knotengrad und auf dem als bekannt vorausgesetzten Pathloss-Exponenten eine Berechnungsvorschrift für die Anpassung der aktuellen Sendeleistung hergeleitet. Dabei liegt die Annahme einer Gleichverteilung der Stationen in der lokalen Umgebung zugrunde.

Die Arbeitsweise des Algorithmus *MobileGrid* [LL02] ist der von LINT sehr ähnlich. Hier wird die Sendeleistung jedoch nicht unmittelbar aufgrund der Nachbarzahl geregelt, sondern aufgrund des *Contention-Index* (CI), der an jeder Station gemessen wird und der angibt, wie viele Stationen sich in einem Zeitfenster das Übertragungsmedium teilen. Es werden also bewusst nur aktive Stationen berücksichtigt mit der Begründung, dass die Zahl der Stationen begrenzt werden muss, die um den Medienzugang konkurrieren, so dass inaktive Stationen keine Rolle spielen. Ebenso wie beim zuvor vorgestellten LINT-Algorithmus liegt die Annahme zugrunde, dass die Sendeleistung auf einen beliebigen Wert in einem kontinuierlichen Bereich eingestellt werden kann, und es wird eine Berechnungsvorschrift zur Anpassung der Sendeleistung im Bedarfsfall angegeben.

Die oben beschriebene Schwäche des Max-Degree-Ansatzes äußert sich in diesem Verfahren darin, dass der implizit angenommene Zusammenhang zwischen dem CI einer Station (also der Anzahl anderer Stationen, die als Absender von Daten- oder Kontrollpaketen wahrgenommen werden können) und ihrer eigenen Sendeleistung bei näherer Betrachtung gar nicht klar ist. Ein unmittelbarer Einfluss der eigenen Sendeleistung auf den CI besteht offensichtlich nicht, da der CI an einer bestimmten Station ausschließlich von den Sendeleistungen anderer Stationen abhängt. Es besteht jedoch ein mittelbarer Einfluss, denn durch das Anpassen der eigenen Sendeleistung verändert sich der CI der Nachbarstationen, so dass komplizierte Wechselwirkungen stattfinden: Eine Station mit zu hohem CI verringert ihre Sendeleistung und damit auch den CI anderer Stationen, die daraufhin ihre Sendeleistungen erhöhen, wodurch sich der CI wieder anderer Stationen erhöht, die ihre Sendeleistungen deshalb wiederum verringern.

Der Vollständigkeit halber ist hier noch der in [KKW03] vorgestellte *Local-Mean-Algorithm*¹ (LMA) zu erwähnen, der genau wie der LINT Algorithmus arbeitet. Zur Erkennung bidirektionaler Links wird vorgeschlagen, dass Stationen in regelmäßigen Abständen Beacons als Broadcast-Rahmen verschicken, die von anderen Stationen explizit quittiert werden, was zu einem hohen Protokolloverhead führen dürfte. Weiterhin wird in dieser Arbeit der *Local-Mean-of-Neighbours-Algorithmus* (LMN) vorgestellt, der analog zu LMA arbeitet, die Entscheidung über Erhöhen oder Verringern der Sendeleistung einer Station jedoch nicht anhand der Nachbarzahl dieser Station alleine trifft, sondern anhand der durchschnittlichen Nachbarzahl der Station und ihrer Nachbarn. Dieser Ansatz wird allerdings nicht motiviert, es wird lediglich ohne Begründung festgestellt, dass Partitionierungen des Netzes bei Verwendung von LMN seltener und weniger schwerwiegend sind.

4.2.1.2 Nearest-Neighbours-Topologien

Die im vorangegangenen Abschnitt vorgestellte Klasse der Max-Degree-Topologien ist das idealisierte Resultat einer Reihe von Verfahren zur dezentralen Sendeleistungsregelung, deren Ziel es ist, die Nachbarzahl jeder Station innerhalb eines gewissen Bereichs zu halten. Als

¹Aus der Veröffentlichung geht nicht hervor, wie diese Bezeichnung zustande kommt; eine Durchschnittsbildung wird in dem Verfahren jedenfalls nicht vorgenommen.

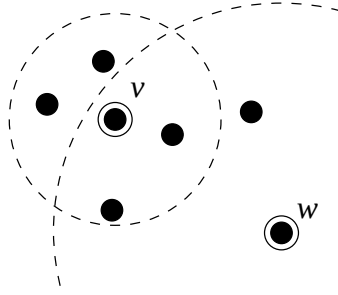


Abbildung 4.3: Die 4-Nearest-Neighbours-Mengen zweier Knoten

Nachteil dieser Topologiekategorie wurde festgestellt, dass Stationen nicht zwingend mit ihren nächsten Nachbarn verbunden sind, sondern teilweise Links zu weiter entfernten Stationen aufbauen. Der Nearest-Neighbours-Graph (siehe z.B. [EPY97]), der häufig im Zusammenhang mit Clusteringverfahren und geometrischen Algorithmen anzutreffen ist, ist dagegen explizit so definiert, dass Verbindungen nur zwischen nächsten Nachbarn bestehen. Ein k -Nearest-Neighbours-Graph einer Punktmenge V ist also definiert als $G = (V, E)$ mit $(v, w) \in E$ genau dann, wenn $w \in N_k(v)$, wobei $N_k(v) \subseteq V \setminus \{v\}$ eine k -Nearest-Neighbours-Menge des Knotens $v \in V$ ist, die genau die k Knoten enthält, die v am nächsten sind. Formal lassen sich Nearest-Neighbours-Mengen rekursiv wie folgt definieren:

$$\begin{aligned} N_0(v) &= \emptyset \\ N_k(v) &= \{w \in V : d_{vw} \leq d_{vu}, \forall u \in V \setminus (N_{k-1}(v) \cup \{v\})\} \end{aligned}$$

Abbildung 4.3 veranschaulicht die 4-Nearest-Neighbours-Mengen zweier Knoten.

Eine k -Nearest-Neighbours-Topologie enthält Links zwischen den Stationen, die im k -Nearest-Neighbours-Graphen adjazent sind. Eine (ungerichtete) Kante zwischen zwei Stationen besteht in der k -Nearest-Neighbours-Topologie also genau dann, wenn eine der beiden Stationen sich unter den k nächsten Nachbarn der anderen Station befindet:

$$(v, w) \in L \Leftrightarrow v \in N_k(w) \vee w \in N_k(v).$$

Aus $v \in N_k(w)$ folgt nicht generell $w \in N_k(v)$, wie Abbildung 4.3 beispielhaft darstellt; die entsprechende 4-Nearest-Neighbours-Topologie enthält also die Kante (v, w) , obwohl $w \notin N_4(v)$. Im Allgemeinen haben in einer k -Nearest-Neighbours-Topologie daher einige Knoten einen höheren Grad als den Mindestgrad von k .

Dies macht den Unterschied zu den Max-Degree-Topologien aus, die im vorangegangenen Abschnitt vorgestellt wurden: Während dort in erster Linie dafür gesorgt wird, dass ein maximaler Knotengrad nicht überschritten wird, ist das vorrangige Ziel hier, dass ein minimaler Knotengrad nicht unterschritten wird. Dadurch werden ungünstige Situationen wie die in Abbildung 4.2 skizzierte verhindert. Erstaunlicherweise ist aber noch kein praktikabler dezentraler Algorithmus zur Sendeleistungsregelung vorgestellt worden, der auf dem Nearest-

Neighbours-Ansatz basiert (mit Ausnahme der im Rahmen der vorliegenden Arbeit entstandenen Veröffentlichungen [GdWMJ03] und [GdWMJ05], siehe Kapitel 7).

Auch der Mutual-Nearest-Neighbours-Ansatz findet sich schon in der Literatur zu Clusteringalgorithmen ([GK78]): Im Gegensatz zu obiger Definition des Nearest-Neighbours-Graphen ist die Kantenmenge E eines *Mutual- k -Nearest-Neighbours-Graphen* dadurch gegeben, dass

$$(v, w) \in L \Leftrightarrow v \in N_k(w) \wedge w \in N_k(v).$$

Der Mutual-4-Nearest-Neighbours-Graph enthält in der in Abbildung 4.3 dargestellten Situation die Kante (v, w) also nicht.

In [EKCD00] wird ein Verfahren vorgestellt, das prinzipiell sowohl Nearest-Neighbours- als auch den Mutual-Nearest-Neighbours-Topologien erzeugen kann. Das vorgestellte Protokoll erfordert eine Synchronisation aller Stationen und einen eigenen Kanal für Signalisierungsnachrichten, so dass ein spezielles Medienzugangsprotokoll in der Lage ist, jeder Station im Signalisierungskanal einen eigenen Zeitslot zu reservieren. Durch die Messung und Mittelung der Signalstärken der empfangenen Pakete, die von allen Stationen periodisch in diesem Kanal mit maximaler Sendeleistung versendet werden, soll eine Station die Dämpfung entlang aller Links abschätzen, an denen sie beteiligt ist, und anhand dieser Werte ihre eigene Sendeleistung für den Datenkanal so regulieren, dass sie gerade ihre k nächsten Nachbarn erreicht. Dieses Verfahren stellt also sehr hohe Anforderungen an die drahtlose Hardware und ist damit gemäß den in Abschnitt 3.3 formulierten Bedingungen nicht praktikabel.

Dem in [BLRS03] vorgestellten *k-neigh protocol*, das auf den energiesparenden Betrieb statischer Netze ausgelegt ist, liegt ebenfalls das Mutual-Nearest-Neighbours-Kriterium zugrunde: In einer initialen Phase berechnet jede Station ihre k -Nearest-Neighbours-Menge durch die Schätzung der Distanzen zu anderen Stationen, die dadurch ermöglicht werden soll, dass alle Stationen Kontrollpakete mit maximaler Leistung versenden. Die berechneten Nachbarmengen werden wieder von allen Stationen mit maximaler Leistung versandt. Die Sendeleistungen werden dann so eingestellt, dass bidirektionale Links zwischen Stationen bestehen, die gegenseitig zu ihren k nächsten Nachbarn gehören. Auch dieses Verfahren wäre durch eine geringfügige Modifikation in der Lage, dem Nearest-Neighbours-Ansatz zu folgen, es ist jedoch aus mehreren Gründen nicht praktikabel. Zunächst ist es für dynamische Netze ungeeignet, da der Overhead der mit maximaler Leistung versendeten Pakete während einer einmaligen Initialisierungsphase tolerierbar sein mag, in dynamischen Netzen müsste dies jedoch regelmäßig geschehen, was aus den in Abschnitt 3.3 beschriebenen Gründen nicht akzeptabel ist. Ein grundsätzlicher Kritikpunkt ist außerdem, dass die Schätzung einer Entfernung basierend auf der Empfangssignalstärke eines einzelnen Pakets sehr ungenau ist.

Eine mit dem Nearest-Neighbours-Ansatz inhärent verbundene Schwäche ist die Tatsache, dass die resultierenden Topologien unter Umständen eine geringere Konnektivität aufweisen, als würden die Stationen maximale Sendeleistungen verwenden. Dies ist in Abbildung 4.4 für $k = 4$ veranschaulicht. Mit steigendem k wird es allerdings immer unwahrscheinlicher, dass ein solcher Fall vorkommt. Außerdem wird in [HM04b] festgestellt, dass die Konnektivität drahtloser Netze durch die Berücksichtigung des Lognormal-Shadowing (siehe Abschnitt

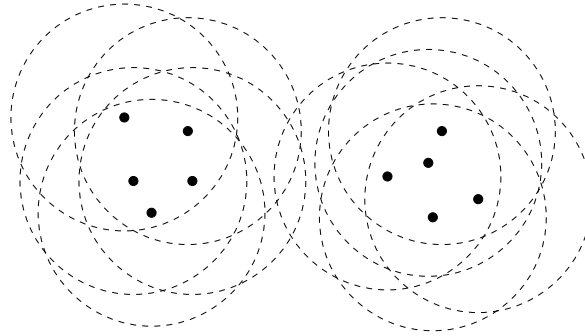


Abbildung 4.4: Schwäche des Nearest-Neighbours-Ansatzes ($k = 4$)

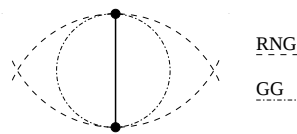


Abbildung 4.5: Bereiche um eine Kante, die in Relative-Neighbourhood- und Gabriel-Graphen keine Knoten enthalten

2.1.3.1) steigt. Auf theoretischer Ebene wird der Zusammenhang zwischen Mindestnachbarzahl und Verbundenheit in [XK04] genau beleuchtet.

4.2.1.3 Minimale Spannbäume

In [LHS03] wird ein Verfahren präsentiert, das minimale Spannbäume in lokalen Umgebungen aufbaut, um möglichst geringe Sendeleistungen zu verwenden, während die Konnektivität des Netzes beibehalten wird. Eine genauere Untersuchung der Eigenschaften solcher Topologien findet sich in [LWS04].

Das in [LHS03] beschriebene Verfahren ist in erster Linie für statische Netze konzipiert und damit nur schwierig in dynamischen Netzen anwendbar, wie zu Beginn des Abschnitts 4.2 ausgeführt. Wie außerdem in Abschnitt 6.5 noch gezeigt wird, weisen minimale Spannbäume wegen ihrer großen Dehnung und Anfälligkeit zur Bildung von Engpässen eine äußerst geringe Transportkapazität auf, weshalb solche Topologien als ungünstig erachtet werden müssen, auch wenn sie aufgrund ihrer minimalen Sendeleistungen und Knotengrade bei Aufrechterhaltung der Konnektivität zunächst interessant erscheinen mögen.

4.2.1.4 Relative-Neighbourhood-Topologien

Der *Relative-Neighbourhood-Graph* (RNG) einer Punktmenge V in einer Ebene ist nach [Tou80] gegeben durch $G = (V, E)$ mit

$$(v, w) \in E \Leftrightarrow d_{vw} \leq \max\{d_{vu}, d_{wu}\}, \quad \forall u \in V \setminus \{v, w\}.$$

Falls d die euklidische Distanz zwischen zwei Knoten bezeichnet, so enthält der RNG eine Kante genau dann, wenn ein bestimmter Bereich, der durch die Positionen der beiden Knoten festgelegt ist, keine weiteren Knoten enthält. Dieser Bereich ist in Abbildung 4.5 dargestellt.

Ein Verfahren zur verteilten Erzeugung des RNG wird in [BJ02] präsentiert, allerdings handelt es sich hierbei um ein geometrisches Verfahren, das ausschließlich dann anwendbar ist, wenn das Distanzmaß tatsächlich die euklidische Distanz der Stationspositionen in einer Ebene ist. Zur Umsetzung dieses Verfahrens müssen die Stationen in der Lage sein, den Winkel zwischen zwei ihrer Links bestimmen zu können.

In [WZ04] wird ein verteiltes Verfahren präsentiert, das den RNG gemäß obiger Definition für ein beliebiges Distanzmaß erzeugt. Die Stationen müssen dazu ihre Links nach ihrer Distanz sortieren. Lediglich dann, wenn die Distanzen nicht paarweise verschieden sind, kann sich die erzeugte Topologie vom RNG unterscheiden, da Links mit gleicher Distanz in willkürlicher Reihenfolge abgearbeitet werden. Wie die Autoren feststellen, können hier aus praktischen Gesichtspunkten Linkmetriken wie die Paketverlustrate verwendet werden. Dieses Verfahren ist allerdings nur für statische Netze konzipiert und kann aus den in Abschnitt 3.3 dargelegten Gründen nicht in dynamischen Netzen betrieben werden, weil dazu regelmäßig mit maximaler Leistung versendete Kontrollnachrichten notwendig wären.

4.2.1.5 Gabriel-Graphen

Der Gabriel-Graph [GS69] einer Punktmenge V ist gegeben durch $G = (V, E)$ mit

$$(v, w) \in E \Leftrightarrow d_{vw} \leq \sqrt{d_{vu}^2 + d_{wu}^2}, \quad \forall w \in V \setminus \{v, w\}.$$

Für die euklidische Distanz zwischen Punkten in einer Ebene bedeutet dies anschaulich, dass zwei Knoten genau dann verbunden sind, wenn in dem Kreis mit dem Durchmesser, der der Distanz zwischen den beiden Knoten entspricht und auf dessen Außenlinie beide Knoten liegen, keine anderen Knoten enthalten sind (siehe Abbildung 4.5). Die Kantenmenge des Gabriel-Graphs ist also eine Obermenge der RNG-Kantenmenge.

In [BDEK02] wird gezeigt, dass der Dehnungsfaktor des Gabriel-Graphen in $O(\sqrt{|V|})$ ist. Im Gegensatz dazu ist der Dehnungsfaktor des RNG in $O(|V|)$. (Zur Definition des Dehnungsfaktors siehe Abschnitt 2.4.2.)

Bisher wurde kein verteiltes Verfahren zur Sendeleistungsregelung vorgestellt, dem dieser Graphyp zugrunde liegt, allerdings wird er an verschiedenen Stellen im Zusammenhang mit der Topologieformung in drahtlosen Netzen erwähnt. Aus den gleichen Gründen wie für RNG-Topologien (siehe Abschnitt 4.2.1.4) wäre ein solches Verfahren für dynamische Netze vermutlich nicht geeignet.

4.2.1.6 Topologien mit vorgegebenem Dehnungsfaktor

In [BvRWZ04] wird zunächst argumentiert, dass viele der veröffentlichten Algorithmen bzw. Strategien zur Sendeleistungsregelung die Interferenz nicht verringerten. Die Interferenz eines Links wird dabei definiert als die Anzahl der Stationen, die sich innerhalb der Sendereichweite mindestens einer der beiden mit dem Link inzidenten Stationen befinden. Mit dieser Definition und den einleitenden Bemerkungen von [BvRWZ04] ist klar, dass Interferenz als die dem Spatial-Reuse entgegengesetzte Größe gesehen wird.

Obwohl in oben angegebener Definition schon kritisiert werden könnte, dass Interferenz über wesentlich größere Entfernungen wirksam sein kann als die Sendereichweite, kann sie in legitimer Weise als Maß dafür dienen, wie viele Stationen blockiert sind, wenn ein bestimmter Link aktiv ist. Die Autoren erweitern den Interferenz-Begriff dann auf ein gesamtes Netz, indem sie die Interferenz im Netz als die maximale Link-Interferenz definieren. Diese Definition reflektiert also die Bedingungen an einem Engpass des Netzes und lässt damit keine Beurteilung darüber zu, ob und wie sich der Spatial-Reuse im Netz insgesamt durch das Verringern der Sendeleistungen erhöht hat.

Weiter werden in [BvRWZ04] Mechanismen zur Erzeugung von Topologien mit minimaler Interferenz bei vorgegebenem Dehnungsfaktor vorgestellt (siehe Abschnitt 2.4.2). Dies wird durch ein einfaches Greedy-Verfahren realisiert: Nach der Reihenfolge ihrer Interferenz-Werte sortiert werden so lange Links zur Topologie hinzugefügt, bis der gewünschte Dehnungsfaktor erreicht ist.

Die Autoren verweisen darauf, dass die Konstruktion eines Pfads zwischen zwei Stationen v und w mit vorgegebener Dehnung t in einem lokal begrenzten Bereich liegen muss, da der Pfad nur Stationen enthalten kann, die höchstens $t/2 \cdot d_{vw}$ von v oder w entfernt sind. Ein solcher Pfad könnte von v oder w also ohne eine globale Sicht des Netzes konstruiert werden, es muss lediglich Topologieinformation aus diesem begrenzten Bereich gesammelt werden. Allerdings kann dieser Bereich leicht sehr groß werden, entweder wenn v und w nicht allzu nahe bei einander liegen, oder wenn eine größere Dehnung zugelassen wird. Um unabhängig von der anliegenden Last einen bestimmten Dehnungsfaktor zu garantieren, müssten außerdem Pfade für alle möglichen Stationspaare gefunden werden, so dass effektiv doch Informationen über alle möglichen Links im gesamten Netz verteilt werden müssten. Damit ist dieser Ansatz kaum praktikabel.

Die zugehörigen Ergebnisse für Beispielnetze liefern aber einen interessanten Einblick in die Wechselwirkungen zwischen dem Grad der Konkurrenz um den Medienzugang und dem Dehnungsfaktor von Topologien: In allen betrachteten Fällen lässt sich die Interferenz durch eine Erhöhung des Dehnungsfaktors von 2 auf 10 höchstens halbieren. Damit ist fraglich, ob die Minimierung der Interferenz wie in [BvRWZ04] definiert als vorrangiges Ziel überhaupt sinnvoll ist, wenn sie durch eine so drastische Erhöhung des Dehnungsfaktors und damit der Pfadlängen erkauft wird.

4.2.1.7 Energie-effiziente Topologien

In [RM98], [RM99] und [LW01a], [LW01b] wird ausgearbeitet, wie Pfade identifiziert werden können, entlang derer Pakete mit minimalem Gesamtaufwand an Energie versendet werden können. Zu diesem Zweck wird eine Topologie erzeugt, die so beschaffen ist, dass ein direkter Link von einer Station v zu einer Station w nur dann besteht, wenn eine direkte Übertragung von v noch w weniger Energieaufwand erfordert als die Übertragung über eine Zwischenstation u :

$$(v, w) \in E \Leftrightarrow d_{vw} \leq d_{vu} + d_{uw}, \quad \forall u \in V \setminus \{v, w\}.$$

Bezeichnet d die euklidischen Distanzen zwischen den Stationen, so ist obige Dreiecksungleichung immer erfüllt. Die Kosten einer Übertragung (in diesem Zusammenhang die für eine vorgegebene mittlere Empfangsleistung benötigte Sendeleistung) steigen aber polynomial mit der Distanz zwischen Sender und Empfänger (siehe Abschnitt 2.1.3.1).

In den erwähnten Arbeiten werden die Kosten aus der Distanz zwischen Sender und Empfänger geschätzt, so dass der Pathloss-Exponent als bekannt vorausgesetzt werden muss. Damit hat dieser konkrete Ansatz wie auch die geometrischen Ansätze verschiedene Nachteile: Neben der Bedingung, dass die Geräte die Fähigkeit dazu haben müssen, die Distanzen zueinander bestimmen zu können, ist die Güte der Schätzung der Linkkosten angesichts der Schwankungen des tatsächlichen Largescale-Pathloss um einen distanzabhängigen Mittelwert, die insbesondere in urbaner Umgebung beobachtet werden, fraglich.

4.2.2 Geometrische Ansätze

Die im Folgenden vorgestellten Ansätze erfordern, dass die Positionen der Stationen oder daraus ableitbare Information, konkret die Winkel zwischen verschiedenen Links einer Station, bekannt sind. Dies ist mit verschiedenen Nachteilen verbunden.

Zunächst wird zur Beschaffung solcher Information spezielle Hardware benötigt. Je nach Ansatz müssen die Stationen entweder mit Hilfe spezieller Antennen zumindest in der Lage sein, den Winkel zwischen zwei Nachbarn zu bestimmen, oder sie müssen durch eine GPS- oder Galileo-Hardware ihre geographische Position bestimmen und diese ihren Nachbarn mitteilen.

Werden tatsächlich die Stationspositionen in einem euklidischen Raum als Eingabe verarbeitet, so sind andere Distanzmaße als die euklidische Distanz, z.B. die Sendeleistung, die notwendig ist, um eine Verbindung zwischen zwei Stationen aufzubauen, nicht anwendbar. Die euklidische Distanz ist aber nur dort sinnvoll einsetzbar, wo die Korrelation zwischen Distanz und Dämpfung des Funksignals hinreichend groß ist, also in erster Linie in offenen Umgebungen. Sind beispielsweise viele größere Hindernisse vorhanden, kann es leicht vorkommen, dass eine Verbindung zwar nur eine geringe Distanz abdeckt, wegen eines Hindernisses aber trotzdem eine hohe Sendeleistung erfordert oder gar nicht aufgebaut werden kann. In urbaner Umgebung oder innerhalb von Gebäuden haben viele Verbindungen starke Multipfad-Anteile

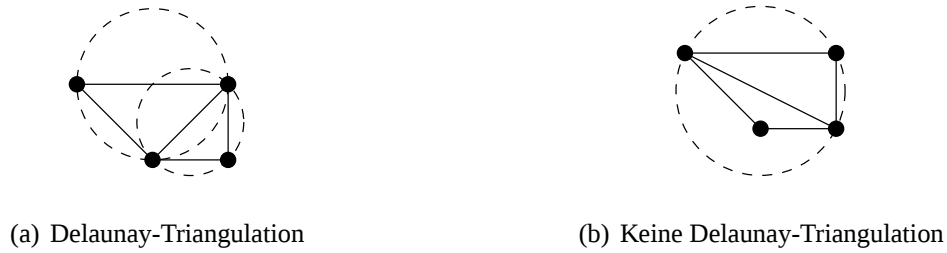


Abbildung 4.6: Beispiel-Triangulationen einer Menge von vier Punkten

und die direkte Sichtlinie ist häufig vernachlässigbar, so dass solche Verfahren dort gar nicht sinnvoll anwendbar sind.

Auch dann, wenn an Stelle der Stationspositionen nur die Winkel zwischen Links als Eingabe dienen, so dass die Verfahren prinzipiell mit anderen Distanzmaßen als der euklidischen Distanz anwendbar wären, bleibt die genannte Einschränkung auf offene Umgebungen bestehen. Die Richtung, aus der ein Signal empfangen wird, lässt sich nämlich nur dann ermitteln, wenn eine wesentliche Komponente des empfangenen Signals entlang der Sichtlinie den Empfänger erreicht. In dicht bebauten oder geschlossenen Umgebungen ist die Wahrscheinlichkeit dagegen hoch, dass diese Komponente von anderen Multipfad-Komponenten überdeckt wird oder gar nicht in nennenswerter Stärke vorhanden ist, weil die Sichtlinie zwischen Sender und Empfänger blockiert ist. Ebenso erfordert die Positionsbestimmung mittels GPS- oder Galileo-Hardware, dass die Geräte sich im Freien befinden.

4.2.2.1 Delaunay-Triangulation

Der älteste Ansatz dieser Art ist vermutlich der *Novel-Topology-Control-Algorithmus* (NTC) [Hu93]. Dieser erzeugt eine *Delaunay-Triangulation* der Stationspositionen: Diese enthält eine Kante zwischen zwei Knoten $v, w \in V$ genau dann, wenn sich deren Voronoi-Zellen berühren, wenn also mindestens ein Punkt $p \in \mathbb{R}^2$ existiert mit

$$d_{pv} = d_{pw} < d_{pu}, \forall u \in V \setminus \{v, w\}.$$

Eine Triangulation bezeichnet allgemein einen aus Dreiecken bestehenden, planaren Graphen. Die Delaunay-Triangulation zeichnet sich dadurch aus, dass jeder Umkreis eines Dreiecks außer den Eckpunkten keine weiteren Knoten enthält. Dies wird in Abbildung 4.6 veranschaulicht.

Der Vorteil dieses Graphen ist die Garantie eines konstanten Dehnungsfaktors. Allerdings sieht der NTC-Algorithmus vor, dass nur bestimmte Links aus der Triangulation übernommen werden, so dass jeder Knoten einen zuvor festgelegten Höchstgrad nicht überschreitet.

Durch diese zusätzliche Bedingung wird (entgegen der Behauptung in [Hu93]) der Vorteil eines solchen Ansatzes, die Garantie maximaler Konnektivität, zunichte gemacht, denn beim

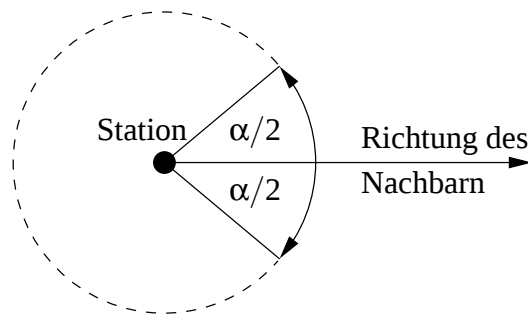


Abbildung 4.7: Veranschaulichung des bei CBTC durch einen Nachbarn abgedeckten Bereichs

Entfernen überzähliger Kanten wird ausschließlich deren Länge betrachtet und nicht ihre Bedeutung für die Konnektivität.

4.2.2.2 Cone-based Topology Control

Der *Cone-based-Topology-Control-Algorithmus* (CBTC) [WLBW01] ähnelt in gewisser Weise dem Nearest-Neighbours-Ansatz. Der entscheidende Unterschied ist, dass die Stationen sich nicht nur mit einer gewissen Anzahl beliebiger Nachbarn verbinden, sondern dass möglichst in jedem Sektor mit dem Winkel α ein Nachbar erreichbar sein soll. Anders ausgedrückt deckt ein Nachbar, dessen relative Position durch den Winkel ρ gegeben ist, den Bereich $[\rho - \alpha/2; \rho + \alpha/2]$ ab (siehe Abbildung 4.7), und jede Station erhöht so lange ihre Sendeleistung, bis ihr gesamter Umkreis abgedeckt ist. Um festzustellen, welche Stationen mit einer bestimmten Sendeleistung erreichbar sind, verschicken die Stationen mit dieser Leistung Broadcast-Rahmen, deren Empfänger mit gleicher Sendeleistung antworten sollen.

Stationen, für die die Forderung nach einem Nachbarn in jedem Sektor mit Winkel α nicht erfüllbar ist, weil sie sich beispielsweise am Rand der Netzes aufhalten, verwenden in dieser Basis-Version des Algorithmus also maximale Sendeleistung. Als ein erster Optimierungsschritt wird daher vorgeschlagen, dass solche Stationen ihre Sendeleistung auf die kleinste Stufe zurückschalten, bei der die zugehörige Abdeckung gleich bleibt.

Ein zweiter Optimierungsschritt beruht auf der in [WLBW01] bewiesenen Tatsache, dass für $\alpha < 2/3\pi$ ein Link nicht zu Erhaltung der Verbundenheit des Netzes beiträgt, wenn nicht beide Stationen diesen Link zur Maximierung der Abdeckung ihres Umkreis benötigen. Deshalb werden Links, die nur von einer Station benötigt werden, beim *Asymmetric-Edge-Removal* entfernt.

Der als *Pairwise-Edge-Removal* bezeichnete dritte Optimierungsschritt entfernt weitere Kanten, die für $\alpha < 5/6\pi$ beweisbar nicht zur Konnektivität des Netzes beitragen: Ein Link zu einem Nachbarn wird von einer Station als *redundant* eingestuft, wenn ein zweiter Link von dieser Station ausgehend zu einem Nachbarn führt, der näher liegt und wenn diese Links einen

Winkel kleiner als $\pi/3$ aufspannen. An dieser Stelle berücksichtigen die Autoren nun auch die Gefahr, dass das Entfernen zu vieler Links zwar nicht die Verbundenheit an sich gefährdet, aber zu langen Umwegen zwischen Quelle-Senke-Paaren und Engpässen in der Netztopologie führen kann. Deshalb sollen nur redundante Links zu den Nachbarn entfernt werden, die weiter entfernt sind als der am weitesten entfernte Nachbar, der durch einen nicht redundanten Link erreicht wird.

Ein Vorteil dieses Verfahrens ist, dass für $\alpha \leq 5/6\pi$ eine maximale Verbundenheit des Netzes beweisbar ist. Wie aber in [WLBW01] auch angemerkt wird, ist ein solcher Ansatz nur anwendbar in einer offenen Umgebung, in der keine größeren Hindernisse existieren.

4.2.3 Andere Ansätze

In diesem Abschnitt werden Verfahren zur Sendeleistungsregelung vorgestellt, die keiner der beiden oben vorgestellten Kategorien angehören.

4.2.3.1 COMPOW und ClusterPow

In [NKSK02] wird der *COMPOW-Algorithmus* beschrieben, der allen Stationen eines Netzes eine einheitliche Sendeleistung zuordnet. Dies wird dadurch motiviert, dass unterschiedliche Sendeleistungen zu unidirektionalen Links führen, die anderen Protokollen Schwierigkeiten bereiten. Hierfür werden Data-Link-Protokolle mit Quittierungsmechanismen und verschiedene Routingprotokolle aufgeführt.

Der COMPOW-Algorithmus setzt voraus, dass auf jeder Sendeleistungsstufe eine eigene Instanz eines proaktiven Routingprotokolls aktiv ist, so dass jede Station darüber informiert ist, welche Ziele erreichbar sind, wenn das Netz mit einer bestimmten Sendeleistung arbeitet. Die Stationen verwenden dann die minimale Sendeleistung, deren zugehörige Routingtabelle maximale Größe hat.

Die Vorteile dieses Verfahrens liegen auf der Hand. Zum einen ist die größtmögliche Konnektivität des Netzes sichergestellt, und zum anderen werden die negativen Auswirkungen unidirektionaler Links minimiert. Letzteres ist jedoch mit Skepsis zu betrachten: Die Klassifizierung und angemessene Behandlung von Links ist vergleichsweise einfach zu realisieren und bringt selbst dann große Vorteile, wenn alle Stationen die gleiche Sendeleistung verwenden, wie in Abschnitt 8.3.1.1 gezeigt wird.

Es ist weiterhin offensichtlich, dass dieses Verfahren mehrere Schwächen hat. Zum einen wird eine große Menge an Kontrollverkehr erzeugt. In vielen Szenarien wird proaktives Routing an sich schon als zu ressourcenaufwändig angesehen, und hier wird insbesondere auch auf hohen Sendeleistungsstufen Routingverkehr erzeugt, selbst wenn dies je nach den Gegebenheiten vermieden werden könnte und sollte. Je mehr Sendeleistungsstufen die Hardware unterstützt, desto unwahrscheinlicher ist es, dass die große Menge an Kontrollverkehr überhaupt noch Ressourcen für den Datenverkehr übrig lässt. Außerdem ist auch unklar, wie dieses Protokoll

arbeiten soll, wenn die Stationen im Netz verschiedene Sendeleistungen unterstützen, was in großen Netzen, in denen Hardware unterschiedlicher Hersteller und unterschiedlichen Alters eingesetzt wird, sehr wahrscheinlich ist.

Das in [KK03] vorgestellte Verfahren *ClusterPow* basiert ebenfalls auf der Idee, auf jeder Sendeleistungsstufe eine eigene Instanz eines proaktiven Routingprotokolls zu betreiben, jedoch dienen die dadurch gesammelten Informationen nicht der Bestimmung einer günstigen einheitlichen Sendeleistung, sondern der Bestimmung eines günstigen Ausgangslinks und der zu verwendenden Sendeleistung für eine konkrete Zielstation. Der Nachteil des hohen Aufkommens an Kontrollverkehr gilt hier also genau wie für [NKSK02].

Bei diesem Verfahren handelt es sich streng genommen also gar nicht um ein Verfahren zur Erzeugung einer Netztopologie; vielmehr wird hier eine Art der Interaktion mit proaktiven Routingprotokollen beschrieben, die es erlaubt, den nächsten Ausgangslink und eine angemessene Sendeleistung für jedes zu versendende Datenpaket zu wählen. Als Fortentwicklung des in diesem Abschnitt eingeführten Verfahrens sollte es hier dennoch kurz erwähnt werden.

4.2.3.2 Berücksichtigung der aktuellen Netzlast

In [PS02a] wird die Idee verfolgt, die Sendeleistungen der Stationen basierend auf der Stärke der in lokalen Umgebungen beobachteten Konkurrenz um den Medienzugang zu regeln. Aus der Feststellung des einfachen Zusammenhangs, dass mit höheren Sendeleistungen auch immer der Durchsatz steigt, leiten die Autoren die Strategie ab, dass Stationen ihre Sendeleistungen erhöhen müssen, wenn die Last steigt. Es handelt sich damit letztlich also um einen Mechanismus zur Senkung des Energiebedarfs: Während die beste Kapazität vermeintlich mit maximalen Sendeleistungen erzielt wird, können bei geringerer Last niedrigere Sendeleistungen verwendet werden.

Hier wird das Potenzial des Spatial-Reuse jedoch unterschätzt, der im Laufe der vorliegenden Arbeit verdeutlicht wird. In [PS02b] erkennen die Autoren sogar, dass mit steigender Stationszahl der Spatial-Reuse immer wichtiger wird, jedoch wird dies abgetan mit der Begründung, dass Szenarien mit hohen Stationszahlen „untypisch“ seien.

Grundsätzlich darf außerdem in Frage gestellt werden, ob eine solch starke Einbeziehung der aktuellen Last sinnvoll ist, wenn die Verkehrsmuster kurzfristigen Schwankungen unterworfen sind. In diesem Fall wird die Topologie laufend angepasst, was wiederum hohe Ansprüche an das Routingprotokoll stellt. Dies sollte gemäß den in Abschnitt 3.3 angestellten Überlegungen jedoch vermieden werden.

4.2.4 Zusammenfassung und Motivation des neuen Verfahrens

In Tabelle 4.1 werden die verschiedenen Ansätze zur Sendeleistungsregelung im Überblick aufgeführt. Wie dort zu sehen ist, erfüllen nur wenige der bisher vorgestellten Verfahren zwei grundlegende Anforderungen, nämlich die Eignung für dynamische Netze und die Unabhän-

Verfahren	Abschnitt	Eignung für dynamische Netze	Benötigt keine geometrische Information
LINT [RRH00], LMA [KKW03]	4.2.1.1	✓	✓
MobileGrid [LL02]	4.2.1.1	✓	✓
[EKCD00]	4.2.1.2	✓	✓
k -neigh [BLRS03]	4.2.1.2		✓
[LHS03]	4.2.1.3		✓
[BJ02]	4.2.1.4		
XTC [WZ04]	4.2.1.4		✓
[BvRWZ04]	4.2.1.6		✓
[RM98], [RM99]	4.2.1.7		
NTC [Hu93]	4.2.2.1		
CBTC [WLBW01]	4.2.2.2		
COMPOW, ClusterPow [NKSK02], [KK03]	4.2.3.1	✓	✓
[PS02a]	4.2.3.2	✓	✓

Tabelle 4.1: Übersicht über Verfahren zur Sendeleistungsregelung

gigkeit von der Verfügbarkeit geometrischer Information, und auch diese sind nur bedingt praktikabel. Die Begründungen werden hier kurz zusammengefasst.

Wie in Abschnitt 4.2.1.1 ausführlich erläutert wurde, ist der Max-Degree-Ansatz nachteilig, weil ungünstig positionierte Stationen eine vorgegebene Mindestzahl von Links gar nicht erreichen können, oder weil sie dafür Links zu unnötig weit entfernten Station aufbauen müssen. Dies wird auch durch die in Kapitel 6 durchgeführten Untersuchungen bestätigt.

Diesen Nachteil hat der Nearest-Neighbours-Ansatz nicht, dessen vorrangiges Ziel die Etablierung einer Mindestzahl von Links für jede Station ist (siehe Abschnitt 4.2.1.2). Das in [EKCD00] vorgestellte Verfahren stellt jedoch unrealistische Forderungen an die verwendete Hardware, und das in [BLRS03] beschriebene Verfahren ist insbesondere für dynamische Netze ungeeignet, weil die Stationen Kontrollnachrichten mit maximaler Leistung senden müssen, um Kenntnis über mögliche Links zu gewinnen.

Wie in Abschnitt 4.2.3.1 erläutert, vervielfachen COMPOW und ClusterPow den Overhead des proaktiven Routings und führen so bei höherer Stationszahl zu einer Überlastung des Netzes mit Kontrollverkehr. Das in [PS02a] beschriebene Verfahren ist ungeeignet, weil es auf einer zweifelhaften Unterschätzung der Möglichkeiten des Spatial-Reuse basiert (siehe Abschnitt 4.2.3.2).

Obwohl der Nearest-Neighbours-Ansatz also vielversprechend für ein praktikables verteiltes Verfahren zur Sendeleistungsregelung ist, ist ein solches in der Forschungslandschaft bisher nicht zu finden (außer in den im Rahmen der vorliegenden Arbeit entstandenen Veröffentli-

chungen [GdWMJ03] und [GdWMJ05]). Nachdem in Kapitel 6 zunächst gezeigt wird, dass Nearest-Neighbours-Topologien tatsächlich vorteilhaft gegenüber anderen Topologien sind, wird in Kapitel 7 daher ein neues verteiltes Verfahren zur Sendeleistungsregelung vorgestellt, das auf dem Nearest-Neighbours-Ansatz basiert.

5 Verfahren zur Kapazitätsabschätzung von Topologien

In diesem Kapitel werden die Verfahren beschrieben, die im folgenden Kapitel 6 zur Messung der Kapazität von Netztopologien dienen. Zunächst wird in Abschnitt 5.1 das Netzmodell erläutert, auf dem diese Verfahren basieren. Abschnitt 5.2 beschreibt die Abschätzung der Broadcast-Kapazität, und in Abschnitt 5.3 wird ein neues Verfahren zur Abschätzung der Transportkapazität einer Netztopologie für ein gegebenes Verkehrsmuster eingeführt.

5.1 Netzmodell

Die Annahmen, die den in diesem Kapitel beschriebenen Verfahren zugrunde liegen, sind im Folgenden aufgeführt. Die globale Synchronisation der Stationen dient zunächst der besseren Greifbarkeit des Problems der Kapazitätsbestimmung. Da die Qualität der Topologien möglichst unabhängig von konkreten Protokollen zur Regelung des Medienzugangs und zur Flusskontrolle bestimmt werden soll, wird weiter angenommen, dass diese Protokolle „perfekt“ arbeiten. Die Berücksichtigung eines Fairness-Kriteriums ist schließlich deshalb nötig, weil der größtmögliche Gesamtdurchsatz häufig durch eine Aufteilung der Bandbreite erreicht wird, die nicht wünschenswert ist, da sie einzelne Instanzen gar nicht berücksichtigt, wie auch das in Abschnitt 2.6 präsentierte Beispiel zeigt.

1. **TDMA-Schema mit Zeitslots konstanter Länge:** Die Stationen sind synchronisiert und halten einen festen, sich periodisch wiederholenden Zeitplan von Übertragungen ein (siehe Abschnitt 2.4.3). Die Dauer aller Zeitslots ist gleich, und in jedem dieser Slots führt eine Station höchstens eine Übertragung durch.

Hier wird also die Ressource „Zeit“ zwischen verschiedenen Entitäten aufgeteilt. Im Broadcast-Fall wird jeder Station die gleiche Sendezeit zugeteilt. Im Fall von Unicast-Datenströmen werden die Sendezeiten max-min-fair zwischen den Strömen aufgeteilt, wobei an jeder Station entlang der Route eines Stroms die gleiche Anzahl von Zeitslots und damit die gleiche Übertragungsdauer für diesen Strom zur Verfügung steht.

Wie genau die Ressourcen „Zeit“ und „Bandbreite“ zueinander in Beziehung stehen, hängt unter anderem von der Länge der zu versendenden Datenrahmen ab. Ist diese

konstant, kann die Länge der TDMA-Slots darauf abgestimmt werden, so dass eine optimale Ausnutzung der Ressourcen stattfindet. Bei variabler Rahmenlänge gehen Ressourcen dadurch verloren, dass Übertragungen die Dauer eines Zeitslots nicht ausnutzen können. Falls die Dauer eines Zeitslots kürzer als die maximale Rahmenlänge ist, entsteht zusätzlicher Fragmentierungs-Overhead.

In der vorliegenden Arbeit wird zur Vereinfachung angenommen, dass alle Pakete die gleiche Länge haben und dass die Stationen mit einer festen Datenrate senden. Damit hängt die Bandbreite, die einer Station oder eine Strom zur Verfügung steht, linear mit dem Quotienten aus der Anzahl der dieser Entität zugeordneten Zeitslots und der Gesamtzahl der Zeitslots im TDMA-Frame zusammen.

2. **Kollisionsfreiheit:** Es wird ein Medienzugang modelliert, der Kollisionen effektiv verhindert; Übertragungen, die sich gegenseitig stören können, finden also nicht im gleichen Zeitslot statt.
3. **Saturierte Sender:** Die Sender haben zu jeder Zeit beliebig viele Daten zu versenden.
4. **Keine Pufferüberläufe:** Im Unicast-Fall ist das Transportprotokoll dazu in der Lage, die Sendefenster basierend auf den verfügbaren Ressourcen optimal zu wählen, so dass keine Pufferüberläufe vorkommen.
5. **Annähernd max-min-faire Ressourcenverteilung:** Im Unicast-Fall werden die Sendefenster außerdem so gewählt, dass die Aufteilung der Ressourcen zwischen den Datenströmen nahezu max-min-fair ist (siehe Abschnitt 2.6). Tatsächlich kann die Aufteilung um höchstens einen Zeitslot pro TDMA-Frame von der max-min-fairen Aufteilung abweichen. Wie in Abschnitt 5.3 genauer erläutert wird, wäre eine perfekt max-min-faire Aufteilung deutlich schwieriger zu realisieren, als dass der Aufwand angesichts der für längere TDMA-Frames vergleichsweise kleinen Abweichung lohnenswert wäre.

5.2 Abschätzung der Broadcast-Kapazität

In Abschnitt 2.4.4 ist die Broadcast-Kapazität eines Netzes als der Durchsatz definiert worden, der sich durch das Fluten (siehe Abschnitt 2.3.1) erzielen lässt. Jedes Paket wird dabei von jeder Station genau einmal als Broadcast-Rahmen an alle Nachbarn verschickt, so dass die Anzahl der Pakete, die jede Station verschickt, gleich ist, unabhängig davon, welche Stationen die Pakete eigentlich generieren. Unter der Annahme von Kollisionsfreiheit wird der Durchsatz in diesem Zusammenhang daher als die Anzahl versendeter Datenpakete pro Zeiteinheit geteilt durch die Stationszahl definiert.

Eine Steigerung des Spatial-Reuse durch das Reduzieren von Sendeleistungen wirkt sich auf die Broadcast-Kapazität grundsätzlich positiv aus, sofern das Netz seine Verbundenheit beibehält. Es erscheint auf den ersten Blick vielleicht widersinnig, dass ein vollständig verbunde-

nes Netz die geringste Broadcast-Kapazität aufweist, da ein neu generiertes Paket beim ersten Versand schon alle anderen Stationen erreicht. Da das Fluten jedoch für den allgemeinen Fall eines nicht zwingend vollständig verbundenen Netzes konzipiert ist und ohne Kenntnis der Netztopologie arbeitet, wird dieses Paket trotzdem von jeder einzelnen Station nochmals übertragen, und jedesmal ist das Medium für alle Stationen im Netz belegt. Aus diesem Grund bietet es sich an, die Broadcast-Kapazität nicht als den erzielbaren Durchsatz selbst zu messen, sondern als Faktor, um den dieser Durchsatz gegenüber einem vollständig verbundenen Netz steigt.

In Abschnitt 2.4.3 ist bereits erläutert worden, dass eine Abschätzung der Broadcast-Kapazität durch ein Distance-2-Vertex-Colouring des Topologiegraphen nur dann sinnvoll ist, wenn alle Links bidirektional sind. Außerdem lässt dies keine allgemeinere Interferenzmodellierung zu. Im gleichen Abschnitt wurde als Lösung das Konzept des Konfliktgraphen präsentiert, der sich in diesem Zusammenhang durch eine Erweiterung der Kantenmenge des Topologiegraphen ergibt. Sind zwei Station im Konfliktgraph adjazent, so bedeutet dies, dass sie nicht gleichzeitig senden dürfen.

Sei S die Menge der Stationen und $L \subseteq S^2$ die Menge der Links. Stört das Signal eines Senders v den Empfang des Signals eines zweiten Senders w (an einem Empfänger u), dann sei $(v, w) \in I$. Damit ergibt sich der Konfliktgraph $G_C = (S, L_C)$ durch

$$L_C = \{(v, w) \in S^2 \mid (v, w) \in L \vee \exists u \in S \setminus \{v, w\} : ((v, u) \in L \wedge (w, u) \in I)\}.$$

Die spezielle Interferenzmodellierung, die das Distance-2-Vertex-Colouring impliziert, ist durch $I = L$ gegeben.

Die chromatische Zahl $\gamma(G_C)$ des Graphen G_C ist die Mindestzahl benötigter Farben für ein Distance-1-Vertex-Colouring des Graphen (siehe Abschnitt 2.4.3). Sie ist also die Mindestzahl von Zeitslots, innerhalb derer jede Station einen Broadcast senden kann, den alle Nachbarn empfangen können. In einem vollständig verbundenen Netz ist $\gamma(G_C) = |S|$, da keine gleichzeitigen Übertragungen stattfinden können. Deshalb gibt der Faktor $|S|/\gamma(G_C)$ die Steigerung der Broadcast-Kapazität im Vergleich zu einem vollständig verbundenen Netz an.

5.3 Der Zwei-Phasen-Färbalgorithmus zur Abschätzung der Transportkapazität

Gemäß der Definition aus Abschnitt 2.4.4 ist die Transportkapazität der maximale erzielbare Gesamtdurchsatz von Unicast-Datenströmen. Wie in Abschnitt 4.1.2.2 dargelegt, sind die bisher entwickelten Verfahren zur Bestimmung der Transportkapazität gegebener Topologien zu rechenaufwändig, um auf Netze mit einer höheren Anzahl an Stationen angewandt werden zu können.

In [GdWFM04] wird die Transportkapazität von Netztopologien durch Simulationen gemessen, in denen drahtlose Multihop-Netze auf sehr abstrakter Ebene modelliert werden, wobei es

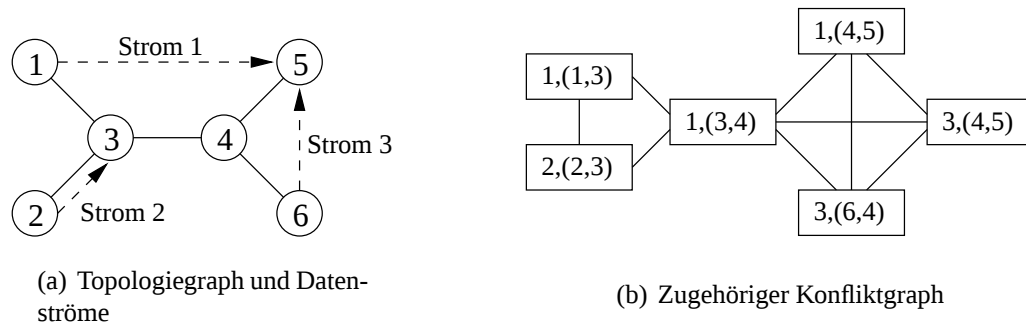


Abbildung 5.1: Beispiel für einen Topologiegraphen und Datenströme und den daraus abgeleiteten Konfliktgraphen

in erster Linie um die Single-Channel-Eigenschaft dieser Netze geht, die durch das Protokoll-Interferenzmodell umgesetzt wird (siehe Abschnitt 2.4.1). Jedoch ist dieses Verfahren immer noch vergleichsweise aufwändig und mit einigen Nachteilen verbunden. Insbesondere ist es sehr schwierig, die verfügbare Bandbreite fair zwischen den Strömen aufzuteilen.

Zur effizienten Abschätzung der Transportkapazität wurde daher der in [Pes04], [PGdWM04] vorgestellte *Zwei-Phasen-Färbalgorithmus* entwickelt. Dieser Algorithmus teilt die verfügbaren Ressourcen möglichst erschöpfend und nahezu max-min-fair auf eine Menge vorgegebener Datenströme auf. Die verwendeten Routen sind ebenfalls Eingabeparameter für den Algorithmus.

Das Netzmodell ist dem in [GdWFM04] verwendeten Netzmodell sehr ähnlich, und deshalb ist es wenig überraschend, dass auch die Ergebnisse eine starke Übereinstimmung zeigen. Während in [GdWFM04] aber noch konkret die Übertragung und Pufferung von Datenpaketen modelliert wird, handelt es sich bei dem im Folgenden präsentierten Verfahren um einen reinen Graphfärbalgorithmus, aus dessen Ausgabe sich ein Übertragungsschedule ableiten lässt.

5.3.1 Vorgehensweise des Algorithmus

Der Algorithmus färbt die Knoten des Konfliktgraphen (siehe Abschnitt 2.4.3), die jeweils eine Übertragung eines Datenstroms repräsentieren. Da ein Link im Allgemeinen von mehreren Datenströmen verwendet wird, können verschiedene Knoten des Konfliktgraphen mit dem gleichen physikalischen Link assoziiert sein. Formal wird ein Knoten des Konfliktgraphen also identifiziert durch ein Tupel (f, l) , wobei f ein Datenstrom ist und $l \in L$ ein Link, der von diesem Datenstrom verwendet wird. Dies wird in Abbildung 5.1 veranschaulicht, wobei die in Abbildung 5.1(b) dargestellte Kantenmenge auf der (für Single-Channel-Netze unrealistischen) Annahme basiert, dass ausschließlich adjazente Links miteinander interferieren.

Eine Farbe entspricht einem Zeitslot im TDMA-Schema. Da der Algorithmus vergleichsweise

Knoten	Phase 1		Phase 2	
	Iteration 1	Iteration 2	Iteration 1	Iteration 2
1, (1,3)	rot	grün	–	–
2, (2,3)	orange	blau	weiß	schwarz
1, (3,4)	gelb	violett	–	–
1, (4,5)	rot	blau	–	–
3, (4,5)	orange	schwarz	–	–
3, (6,4)	grün	weiß	–	–

Tabelle 5.1: Beispielfärbung des Konfliktgraphen aus Abbildung 5.1(b)

lange TDMA-Frames erzeugt und ein Sender im Regelfall mehrere Pakete pro TDMA-Frame verschickt, wird jedem Knoten des Konfliktgraphen nicht nur eine Farbe, sondern eine Menge von Farben zugeordnet. Kollisionsfreiheit ist dann und nur dann gegeben, wenn die Schnittmenge aller Farben, die zwei benachbarten Knoten zugeordnet sind, leer ist. Da ein Transportprotokoll modelliert wird, das die Datenrate so zu wählen vermag, dass Paketverwerfungen nicht vorkommen, werden alle Knoten des Konfliktgraphen, die zu einem bestimmten Strom gehören, mit der gleichen Anzahl an Farben gefärbt.

Der Algorithmus ist in Abbildung 5.2 in Pseudocode-Notation dargestellt. Als Eingabe erhält er den Konfliktgraphen sowie die Zuordnung von Knoten des Konfliktgraphen zu Datenströmen, die hier durch natürliche Zahlen repräsentiert sind. Beide ergeben sich aus der Netztopologie, dem Interferenzmodell und den zuvor festgelegten Routen, entlang derer die Datenströme transportiert werden (siehe Abbildung 5.1). Ein weiterer Eingabeparameter, der den Zeitpunkt des Übergangs von der ersten in die zweite Phase bestimmt, ist der weiter unten erläuterte Stagnations-Schwellwert.

5.3.2 Phase 1: Ressourcenverteilung unter Einhaltung strikter Fairness

Wie der Name des Algorithmus andeutet, werden die Knoten des Konfliktgraphen in zwei Phasen gefärbt. In jeder Iteration der ersten Phase wird jedem Knoten eine zusätzliche Farbe zugeordnet. Wenn möglich, wird dafür eine Farbe verwendet, die bereits einem anderen Knoten zugeordnet ist, ansonsten wird eine neue Farbe eingeführt. Da das Einführen einer neuen Farbe eine Vergrößerung des TDMA-Frames impliziert, wird in der ersten Phase also allen Datenströmen ein möglichst hoher Durchsatz unter Einhaltung strikter Fairness zugeordnet.

Der in Abbildung 5.1(b) dargestellte Konfliktgraph könnte beispielsweise in der ersten Iteration mit den vier Farben Rot, Orange, Gelb und Grün gefärbt werden (siehe Tabelle 5.1). In jeder nachfolgenden Iteration dieser Phase müssen wieder vier neue Farben eingeführt werden, weil die größte Clique gerade die vier Knoten $(1, (3,4))$, $(1, (4,5))$, $(3, (4,5))$ und $(3, (6,4))$ umfasst. (Ein *Clique* ist eine Menge von Knoten, die einen vollständig verbundenen Teilgraphen induziert.) Deshalb kommen in der zweiten Iteration der ersten Phase vier

Abbildung 5.2 : Der Zwei-Phasen-Färbealgorithmus

Parameter : (Ungerichteter) Konfliktgraph $G = (V, E)$
 Zuordnung von Übertragungen zu Datenströmen $m : V \rightarrow \mathbb{N}$
 Stagnations-Schwellwert $t \in \mathbb{N}$

```

forall  $v \in V$  do
     $C_v \leftarrow \emptyset$ 
// Erste Phase:
 $n_{\text{old}} \leftarrow \infty$ 
while  $t > 0$  do
     $n_{\text{new}} \leftarrow 0$ 
    forall  $v \in V$  do
         $I_v \leftarrow \{v' \in V : v' = v \vee (v, v') \in E\}$ 
         $P \leftarrow \bigcup_{v' \in V} C_{v'} \setminus \bigcup_{v' \in I_v} C_{v'}$ 
        if  $P \neq \emptyset$  then
            // Es kann eine bereits eingeführte Farbe wiederverwendet werden.
             $c \leftarrow$  beliebiges Element aus  $P$ 
        else
             $c \leftarrow |\bigcup_{v' \in V} C_{v'}|$ 
             $n_{\text{new}} \leftarrow n_{\text{new}} + 1$ 
         $C_v \leftarrow C_v \cup \{c\}$ 
    if  $n_{\text{new}} \geq n_{\text{old}}$  then
         $t \leftarrow t - 1$ 
     $n_{\text{old}} \leftarrow n_{\text{new}}$ 
// Zweite Phase:
 $F \leftarrow \bigcup_{v \in V} \{m(v)\}$ 
while  $F \neq \emptyset$  do
    forall  $f \in F$  do
         $V_f \leftarrow \{v \in V : m(v) = f\}$ 
         $S \leftarrow \emptyset$ 
        forall  $v \in V_f$  do
             $P \leftarrow \bigcup_{v' \in V} C_{v'} \setminus \bigcup_{v' \in I_v} C_{v'}$ 
            if  $|P| > 0$  then
                 $C_v \leftarrow C_v \cup \{\text{beliebiges Element aus } P\}$ 
                 $S \leftarrow S \cup \{v\}$ 
        if  $|S| < |V_f|$  then
            // Nicht alle Übertragungen konnten gefärbt werden.
             $F \leftarrow F \setminus \{f\}$ 
            forall  $v \in S$  do
                 $C_v \leftarrow C_v \setminus \{\text{beliebiges Element aus } C_v\}$ 
return  $\{(v, C_v) : v \in V\}$ 
    
```

Slot 1 (orange)	Slot 2 (weiß)	Slot 3 (gelb)	Slot 4 (grün)	Slot 5 (schwarz)	Slot 6 (blau)	Slot 7 (rot)	Slot 8 (violett)
2, (2, 3) 3, (4, 5)	2, (2, 3) 3, (6, 4)	1, (3, 4)	1, (1, 3) 3, (6, 4)	2, (2, 3) 3, (4, 5)	2, (2, 3) 1, (4, 5)	1, (1, 3) 1, (4, 5)	1, (3, 4)

Tabelle 5.2: Ein aus Tabelle 5.1 abgeleitetes TDMA-Schema

neue Farben hinzu: Blau, Violett, Schwarz und Weiß.

Im allgemeinen Fall wird die Anzahl der neu eingeführten Farben zunächst mit jeder Iteration kleiner, da immer mehr andere Farben schon zur Verfügung stehen. Nach mehreren Iterationen ist jedoch zu beobachten, dass diese Zahl sich nicht mehr nennenswert verändert. Da nämlich ein bestimmter Bereich des Konfliktgraphen einen Engpass darstellt (wie die Clique in obigem Beispiel), werden in jeder Iteration neue Farben benötigt, nur damit die Ströme in diesem Bereich ein weiteres Paket ausliefern können, während in anderen Teilen des Netzes immer mehr Ressourcen ungenutzt bleiben. Die Stagnation der Anzahl neu eingeführter Farben dient daher als Kriterium für den Übergang in die zweite Phase: Der als *Stagnations-Schwellwert* bezeichnete Eingabeparameter bestimmt, nach wie vielen Iterationen, in denen die Zahl neu eingeführter Farben sich im Vergleich zur vorangegangenen Iteration nicht verringert hat, die erste Phase abgeschlossen wird.

Diese konkrete Bedingung dient hier deshalb als Stagnationskriterium, weil in manchen Fällen beobachtet werden kann, dass die Anzahl neu eingeführter Farben von einer Iteration zur nächsten zwar nicht gleich bleibt, aber verschiedene Werte aus einem kleinen Intervall annimmt. Dadurch könnte die erste Phase unnötig lang werden, wenn auf ein Gleichbleiben der Anzahl neu eingeführter Farben überprüft würde, und es kann nicht garantiert werden, dass der Algorithmus überhaupt terminiert.

5.3.3 Phase 2: Water-Filling

In der zweiten Phase wird die Länge des TDMA-Frames beibehalten, und die verbleibenden freien Ressourcen werden unter den Strömen aufgeteilt. Es werden also keine neuen Farben mehr eingeführt. Kann ein Knoten des Konfliktgraphen mit keiner zusätzlichen bereits anderweitig benutzten Farbe mehr gefärbt werden, wird der gesamte Strom *blockiert*, wird also in den nachfolgenden Iterationen nicht mehr berücksichtigt.

Im Beispiel, das in Tabelle 5.1 skizziert ist, beginnt die zweite Phase nach den zwei bereits beschriebenen Iterationen der ersten Phase (der Stagnations-Schwellwert ist also $t = 1$). Schon in der ersten Iteration dieser zweiten Phase müssen die Ströme 1 und 3 blockiert werden, da für die genannte Clique aus vier Knoten keine Farben mehr verfügbar sind. Strom 2 wird hingegen in zwei weiteren Iterationen mit zusätzlichen Farben versehen, bevor auch er blockiert wird und der Algorithmus terminiert. Zur resultierenden Färbung des Konfliktgraphen korrespondiert ein TDMA-Frame aus acht Slots (siehe Tabelle 5.2), wobei die Reihenfolge der Slots beliebig permutiert werden kann. Innerhalb der Zeitdauer eines TDMA-Frames können

im eingeschwungenen Zustand Strom 1 und Strom 3, die am stärksten Engpass des Netzes beteiligt sind, zwei Pakete an ihre jeweiligen Senken ausliefern, während Strom 2 zusätzliche Ressourcen nutzen kann und in der gleichen Zeit 4 Pakete ausliefert.

Im Gegensatz zu dem einfachen Beispiel, das hier zur Veranschaulichung dient, kann es in der zweiten Phase leicht zu Unfairness zwischen den Strömen kommen, weil deren Abarbeitungsreihenfolge bestimmt, welchen Strömen in einer Iteration noch eine zusätzliche Übertragung zugeordnet wird und welchen nicht. Beispielsweise kann die Situation auftreten, dass noch einige, aber nicht alle Knoten gefärbt werden können, die Übertragungen verschiedener Ströme auf dem gleichen Link repräsentieren, wie Knoten $(1, (4, 5))$ und $(3, (4, 5))$ im Konfliktgraph in Abbildung 5.1(b).

In einem solchen Fall wäre eine wirklich max-min-faire Aufteilung nur zu erzielen, indem die Länge des TDMA-Schemas in geeigneter Weise vervielfacht wird. Kann beispielsweise von zwei Strömen nur einer gefärbt werden, wäre das bisher berechnete TDMA-Schema zu duplizieren, wobei in der einen Kopie ein Strom zu färben ist und in der zweiten Kopie der andere Strom. Durch das wiederholte Auftreten solcher Vervielfachungen des TDMA-Schemas kann dieser schnell extrem lang werden. Außerdem ist es im allgemeinen Fall schon aufwändig, überhaupt zu erkennen, dass eine Ungleichbehandlung von Strömen vorliegt und welche Ströme dabei vom gleichen Flaschenhals, der nicht notwendigerweise ein einzelner Link sein muss, betroffen sind. Auf die Erkennung und Auflösung der Ungleichbehandlung von Strömen wird daher verzichtet, denn der Durchsatz von Strömen, die gemäß einer wirklich max-min-fairen Aufteilung den gleichen Durchsatz erzielen müssten, unterscheidet sich höchstens um nur ein Paket pro TDMA-Frame. Dieser Unterschied ist vernachlässigbar, wenn der TDMA-Frame lang genug ist, wenn in der ersten Phase also genügend Iterationen durchlaufen werden.

5.3.4 Färbeheuristik

Der im vorangegangenen Abschnitt beschriebene Algorithmus lässt bestimmte Details offen, die einen wesentlichen Einfluss auf die Güte des Ergebnisses haben. Eines dieser Details ist die Auswahl des nächsten zu färbenden Knotens in der ersten Phase des Algorithmus. Die Implementierung, die den in dieser Arbeit vorgestellten Ergebnissen zugrunde liegt, verwendet dazu die bekannte *Saturation-Largest-First-Heuristik* (SLF) [Bré79]. Nach dieser Heuristik hat der Knoten Vorrang, der die größte Zahl (in der laufenden Iteration) bereits gefärbter Nachbarn hat. Trifft dies auf mehrere Knoten zu, wird unter ihnen derjenige mit dem höchsten Knotengrad gewählt. Sofern auch dieser Knoten nicht eindeutig ist, entscheidet in letzter Instanz der Knotenindex.

Da in der ersten Phase die Stromzugehörigkeit eines Knotens keine Rolle spielt, sondern jedem Knoten genau eine neue Farbe pro Iteration zugewiesen wird, ist diese Heuristik ohne weiteres umsetzbar. In der zweiten Phase werden die Knoten des Konfliktgraphen jedoch nach Stromzugehörigkeit abgearbeitet. Hier muss also entschieden werden, welche Übertragungen eines Datenstroms in ihrer Gesamtheit als nächstes bearbeitet werden sollen. Diese Abarbei-

tungsreihenfolge lehnt sich in der Implementierung, die für die vorliegende Arbeit benutzt wurde, an die SLF-Heuristik an. Anstelle der Anzahl bereits gefärbter Nachbarn und dem Grad eines einzelnen Knotens werden hier die entsprechenden Durchschnittswerte über alle Knoten eines Stroms betrachtet, und in letzter Instanz entscheidet der Stromindex.

Eine Alternative zur Färbung der Knoten in der zweiten Phase wäre es gewesen, die Knoten, die zu nicht blockierten Strömen gehören, genau wie in der ersten Phase unabhängig von ihrer konkreten Stromzugehörigkeit in der von SLF vorgegebenen Reihenfolge abzuarbeiten. Kann ein Knoten dann nicht mehr gefärbt werden, wird der zugehörige Strom genau wie in der bereits vorgestellten Variante des Algorithmus blockiert und jedem Knoten dieses Stroms, der in dieser Iteration bereits gefärbt wurde, eine Farbe wieder entzogen. Die Situation wird jedoch dadurch wesentlich komplizierter, dass sich durch diese frei gewordenen Ressourcen andere Stromblockierungen als fehlerhaft herausstellen können: Die Farben, die nun entfernt wurden, hätten vielleicht auch adjazanten Knoten zugeordnet werden können, die zuvor nicht weiter gefärbt werden konnten und so zur Blockierung ihres Strom geführt haben. Deshalb müssen die Knoten all jener Ströme, auf die das zutreffen könnte, in der laufenden Iteration erneut berücksichtigt werden. In [Pes04] ist auch diese Vorgehensweise verfolgt worden, es hat sich jedoch gezeigt, dass die Ergebnisse sich kaum von denen unterscheiden, die mit der einfacheren Variante erzielt wurden, die die Knoten nach Stromzugehörigkeit gruppiert abarbeitet.

5.3.5 Vergleich mit anderen Verfahren

Das im Vorangegangenen eingeführte Verfahren hebt sich zunächst dadurch von anderen Verfahren zur Bestimmung der Transportkapazität (siehe Abschnitt 4.1.2.2) ab, dass auch Topologien mit einer hohen Anzahl von Stationen in kurzer Zeit bearbeitet werden können. Im folgenden Kapitel werden Durchsatzwerte für 250 Ströme in Netzen mit 1000 Stationen präsentiert; die Färbung eines Konfliktgraphen für ein solches Szenario dauerte mit der heute üblichen PC-Hardware nur wenige Sekunden.

Eine Ursache für diese Verbesserung ist auch, dass die zur Datenübertragung verwendeten Routen Teil der Eingabe des Zwei-Phasen-Färbalgorithmus sind. Dadurch wird das Problem leicht vereinfacht, der erzielbare Durchsatz ist jedoch eingeschränkt auf eine vorgegebene Routingstrategie. Dies kann einerseits von Vorteil sein; so werden in Abschnitt 6.5 Schlussfolgerungen aus Unterschieden im erzielbaren Durchsatz für verschiedene Routingstrategien gezogen. Andererseits ist es mit diesem Verfahren nicht möglich, den unabhängig von der konkreten Routing-Strategie maximal erzielbaren Durchsatz zu bestimmen.

Dazu ist allerdings anzumerken, dass der sparsame Umgang mit den zur Verfügung stehenden Ressourcen notwendig zur Maximierung des Durchsatzes ist, wenn genügend viele Ströme das Netz auslasten. Werden nur wenige Ströme modelliert, so ist es unter Umständen möglich, durch die Wahl längerer Routen ansonsten ungenutzte Ressourcen zu verwenden und damit die Bildung von Bottlenecks zu vermeiden, so dass der Durchsatz entsprechend gesteigert werden kann. Je mehr Ströme aber hinzukommen, desto weniger wahrscheinlich ist

es, dass die Ressourcen entlang alternativer Routen eines Stroms ungenutzt sind. Das Routing entlang eines Umwegs bedeutet dann, dass der zugehörige Strom Ressourcen ineffizient einsetzt, die andere Ströme effizienter nutzen könnten. Dies hat praktisch gesehen zur Folge, dass die Routen bei einer höheren Anzahl von Datenströmen auf kürzeste Pfade gemäß einer geeigneten Routingmetrik beschränkt werden müssen.

6 Auswirkungen der Sendeleistungsregelung auf die Netzkapazität

In diesem Kapitel werden mehrere der Strategien zur Topologieformung verglichen, die in Abschnitt 4.2 eingeführt worden sind. Dazu werden konkrete Netztopologien untersucht, die aus fest vorgegebenen Stationspositionen durch Anwendung der verschiedenen Strategien entstehen. Von besonderem Interesse sind dabei jene Strategien, die, wie in Abschnitt 3.3 gefordert, keine speziellen Anforderungen an die Hardware oder die Einsatzumgebung stellen.

Abschnitt 6.1 beschreibt zunächst die Annahmen, die den in diesem Kapitel beschriebenen Auswertungen zugrunde liegen. In Abschnitt 6.2 ist zusammengefasst, welche Strategien zur Topologieformung betrachtet werden. Außerdem sind dort einige praktische Erwägungen zur konkreten Implementierung verschiedener Strategien beschrieben. Abschnitt 6.3 führt dann einige Metriken ein, anhand derer sich bestimmte Topologieeigenschaften quantifizieren lassen. In Abschnitten 6.4 und 6.5 wird der Einfluss der Topologieformung auf die Broadcast-Kapazität und die Transportkapazität drahtloser Netze untersucht.

6.1 Modellierung

6.1.1 Links

Die Modellierung der Netze erfolgt wie in Abschnitt 2.4.1 beschrieben. Die gegebenen Szenarien bestehen aus einer Menge von Stationen S und einer Distanzfunktion $d : S \times S \rightarrow \mathbb{R}$, die hier die euklidische Distanz zwischen den zuvor festgelegten Positionen der Stationen angibt. Die Topologie ergibt sich durch Festlegung der Sendereichweiten der Stationen $r : S \rightarrow \mathbb{R}$. Ein Link (v, w) von Station v zu Station w besteht genau dann, wenn $d_{vw} \leq r_v$, wenn sich also w innerhalb der Sendereichweite von v befindet.

Bei der Untersuchung der Broadcast-Kapazität wird angenommen, dass Datenübertragungen entlang aller Links, also auch entlang unidirektionaler Links stattfinden, da Broadcast-Rahmen schwierig abzusichern sind. Bei der Untersuchung der Transportkapazität werden hingegen nur bidirektionale Links zugelassen, da Quittungen auf Data-Link-Ebene für Unicast-Rahmen wegen der hohen Fehleranfälligkeit drahtloser Datenübertragungen als unverzichtbar angesehen werden.

Die beschriebene Modellierung der Topologien ist im Einklang mit den in Abschnitt 3.2 erläuterten Überlegungen, dass jeder Station eine bestimmte Sendeleistung zugeordnet wird, die sie für alle Übertragungen unabhängig von deren Empfängern verwendet. Viele der untersuchten Strategien legen jedoch konkret fest, welche Links die Topologie enthalten soll und welche nicht. In diesen Fällen wird als Sendereichweite einer Station das Maximum der Distanzen aller Links dieser Station angenommen. Dadurch kann die resultierende Topologie zusätzliche Links enthalten, die gemäß der ursprünglich betrachteten Strategie nicht vorgesehen sind.

Die maximale Sendereichweite unterliegt hier generell keiner Beschränkung, mit Ausnahme der Max-Degree-Topologien, für die aus den in Abschnitt 6.2 erläuterten Gründen eine maximale Sendereichweite angenommen wird.

6.1.2 Interferenz

Broadcast-Kapazität Bei der Untersuchung der Broadcast-Kapazität dient als Interferenzmodell eine einfache Erweiterung des Protokollmodells (siehe Abschnitt 2.4.1), dass die Möglichkeit individueller Sendereichweiten berücksichtigt. Eine Übertragung von v nach w ist demnach genau dann erfolgreich, wenn w innerhalb der Sendereichweite von v , aber außerhalb der Störreichweite jedes anderen Senders u ist, die als $(1 + \delta) \cdot r_u$ angenommen wird. Der Parameter $\delta \geq 0$ wird auch als *Guard-Zone* bezeichnet und ermöglicht die Modellierung von Störreichweiten, die größer sind als die Sendereichweiten der Stationen. Der Erfolg einer Übertragung von v nach w ist also genau dann gegeben, wenn

$$d_{vw} \leq r_v \quad \text{und} \quad d_{uw} \geq (1 + \delta) \cdot r_u, \quad \forall u \in S' \setminus \{v\},$$

wobei S' die Menge aller Stationen ist, die im betrachteten Zeitslot senden.

Die Auswertung der Broadcast-Kapazität erfolgt für $\delta = 0$ und $\delta = 1$, also für einen sehr optimistischen Wert und einen eher vorsichtigen Wert für die Guard Zone.

Transportkapazität Wie bereits erläutert, werden bei der Auswertung der Transportkapazität nur bidirektionale Links zugelassen, da drahtlose Unicast-Übertragungen wegen deren hoher Fehleranfälligkeit eine Absicherung durch Quittungen auf der Verbindungsschicht erfordern. Das Interferenzmodell wird deshalb entsprechend modifiziert. Für eine erfolgreiche Übertragung ist es daher erforderlich, dass auch der Sender v außerhalb der Störreichweite der anderen Sender liegt, da sonst der erfolgreiche Empfang einer Quittung nicht sichergestellt ist. Ebenso kann dies auch als Modellierung eines CSMA-Mechanismus gesehen werden.

Eine Übertragung von v nach w ist daher genau dann erfolgreich, wenn

$$d_{vw} \leq \min\{r_v, r_w\} \quad \text{und} \quad \min\{d_{uv}, d_{uw}\} \geq (1 + \delta) \cdot r_u, \quad \forall u \in S' \setminus \{v, w\},$$

wobei S' die Menge der Stationen ist, die im gleichen Zeitslot an einer Übertragung (entweder als Sender oder als Empfänger) beteiligt sind.

In den folgenden Auswertungen ist $\delta = 1$. Stationen, die an einer Übertragung beteiligt sind, blockieren damit alle Stationen, die weniger als die doppelte Sendereichweite entfernt sind. Über welche Entfernung das Vorhandensein eines Signals tatsächlich festgestellt werden kann, hängt stark von der verwendeten Hardware und der Umgebung ab, und eine feste „Störreichweite“ ist sowieso nur ein gedankliches Konstrukt, denn die tatsächliche Störreichweite hängt auch von der Entfernung des zweiten Senders zu der Station ab, an der die Signale interferieren. Der Wert $\delta = 1$ kann aber als eher vorsichtig angesehen werden. Je höher δ gewählt wird, desto geringer ist das Potenzial des Spatial-Reuse, so dass die im Folgenden präsentierten Ergebnisse dieses Potenzial keinesfalls stark überschätzen.

6.1.3 Verkehrsmuster

Der modellierte Verkehr besteht bei der Untersuchung der Transportkapazität aus einer Menge von Datenströmen, deren Quellen und Senken zufällig gemäß einer Gleichverteilung gewählt werden (mit der Einschränkung, dass Quelle und Senke eines Stroms verschieden sein müssen).

Die Auswahl der Quelle-Senke-Paare beeinflusst die Ergebnisse relativ stark. Die konkreten Entfernungen der Quelle-Senke-Paare sind nur dann irrelevant, wenn alle Stationen direkt miteinander verbunden sind, da alle Quellen sich das gemeinsame Übertragungsmedium teilen und die Pakete ihre Senken unmittelbar erreichen. Wenn nun geringere Distanzen zwischen Quellen-Senke-Paaren wahrscheinlicher sind als höhere, so kann eine Steigerung des Gesamtdurchsatzes erreicht werden, indem die Sendereichweiten so verringert werden, dass der Spatial-Reuse deutlich gesteigert wird, während immer noch eine einzige Übertragung oder vergleichsweise wenige Übertragungen nötig sind, um ein Paket von einer Quelle zu einer Senke zu transportieren. Sind die Distanzen zwischen den Quellen und den zugehörigen Senken jedoch mit hoher Wahrscheinlichkeit groß, dann resultiert eine Verringerung der Sendereichweiten in einer deutlich größeren Anzahl der zum Transport eines Paket von einer Quelle zu einer Senke benötigten Übertragungen. Sofern also überhaupt eine Steigerung der Kapazität möglich ist, fällt sie in diesem Fall geringer aus.

Die zufällige, gleichverteilte Auswahl von Quelle-Senke-Paaren ist dann gerechtfertigt, wenn die Kommunikation unabhängig von den Aufenthaltsorten der Netzteilnehmer ist, wie es für Anwendungen wie File-Sharing, Chat, oder Online-Spiele durchaus vorstellbar ist. Eine Bevorzugung kürzerer Entfernungen ist in Szenarien möglich, in denen die Netzteilnehmer an verschiedenen Aufgaben mit lokalem Bezug arbeiten, wobei die Notwendigkeit eines Informationsaustauschs umso wahrscheinlicher ist, je näher die Aufgabengebiete beieinander liegen. In diesem Fall ist die Steigerung, die durch die Regelung der Sendeleistungen erzielbar ist, größer als für die gleichverteilte Auswahl, die den im Folgenden präsentierten Ergebnissen zugrunde liegt.

Sich ein Szenario vorzustellen, in dem Datenströme umso wahrscheinlicher sind, je größer die Quelle-Senke-Distanz ist, fällt dagegen schwer. Dort wäre durch die Regelung der Sendeleistungen jedenfalls weniger Transportkapazität hinzuzugewinnen, als es die im Folgenden

präsentierten Ergebnisse versprechen.

6.2 Untersuchte Strategien zur Topologieformung

Die Strategien zur Topologieformung, die in diesem Kapitel verglichen werden, werden im Folgenden zusammen mit einigen praktischen Erwägungen kurz aufgeführt. Für detailliertere Beschreibungen der einzelnen Strategien wird auf die entsprechenden Abschnitte in Kapitel 4 verwiesen. Abbildung 6.1 stellt außerdem beispielhaft die bidirektionalen Links dar, die aus einer vorgegebenen Menge von 1000 Stationspositionen resultieren (siehe Abbildung 6.1(a)).

Werden die Stationspositionen zufällig gemäß einer kontinuierlichen Verteilung gewählt, sind die Distanzen mit an Sicherheit grenzender Wahrscheinlichkeit eindeutig. Da für die Repräsentation der Positionen nur eine endliche Anzahl von Bits zur Verfügung steht, wird bei der Erzeugung der Szenarien mathematisch gesehen zwar eine diskrete Verteilung verwendet, allerdings ist die Wahrscheinlichkeit dafür, dass eine Distanz mehrfach vorkommt, angesichts der Genauigkeit der Fließkommazahlen immer noch vernachlässigbar gering.

Common-Range-Topologien Diese Topologien basieren auf dem einfachsten Konzept der Topologieformung, dass allen Stationen die gleiche Sendereichweite zugeordnet wird.

An diesen Topologien wird der generelle Einfluss der Sendeleistungsregelung über einen großen Parameterbereich untersucht, der bei so geringen Werten beginnt, dass kaum Links vorhanden sind und sich bis zu einem Maximalwert erstreckt, für den der Topologiegraph vollständig verbunden ist, so dass keine Multihop-Kommunikation mehr notwendig ist.

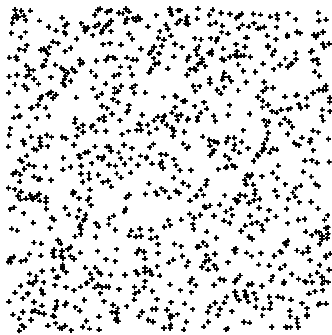
In Abbildung 6.1(b) ist eine Common-Range-Topologie visualisiert. Die Sendereichweite jeder Station beträgt 50m (die Diagonale der quadratischen Simulationsfläche ist 1000m lang).

NNTC-Topologien Diese Topologien werden gemäß der Nearest-Neighbours-Strategie erzeugt (siehe Abschnitt 4.2.1.2). Die Sendereichweiten sind also minimal unter Berücksichtigung der Anforderung, dass jede Station mit ihren k nächsten Nachbarn verbunden sein muss.

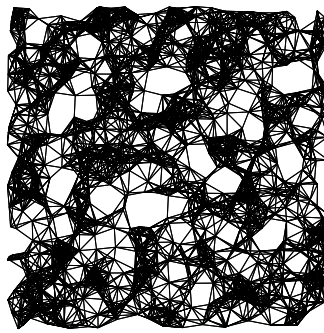
Abbildung 6.1(c) zeigt eine Nearest-Neighbours-Topologie für $k = 6$.

Max-Degree-Topologien Die Konstruktion der Max-Degree-Topologien beruht auf den in Abschnitt 4.2.1.1 vorgestellten Verfahren. Die Idee dieser Verfahren ist, dass eine Station ihre Sendeleistung erhöht, wenn sie weniger als k Nachbarn hat, und ihre Sendeleistung verringert, wenn sie mehr als k Nachbarn hat. Die Sendereichweiten der Stationen werden für diese Topologien also so eingestellt, dass jede Station möglichst viele, aber nicht mehr als k Links hat.

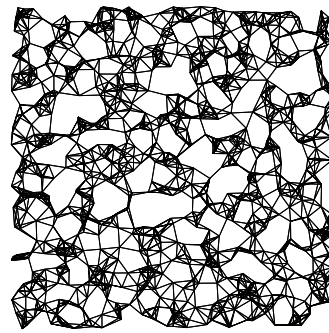
Implementiert wurde dies wie in Abbildung 6.2 in Pseudocode-Notation angegeben: Zunächst wird allen Stationen die Sendereichweite 0 zugeordnet. In einer Schleife wird dann aus allen



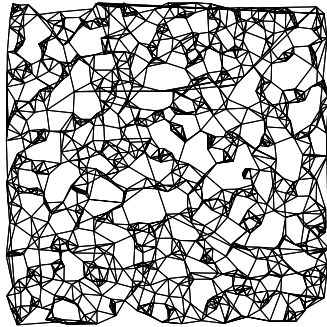
(a) Stationspositionen



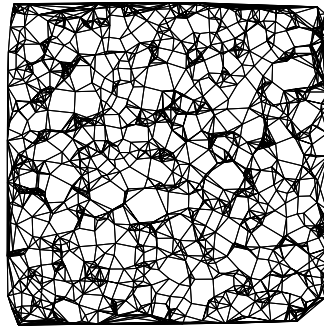
(b) Common Range ($r = 50\text{m}$)



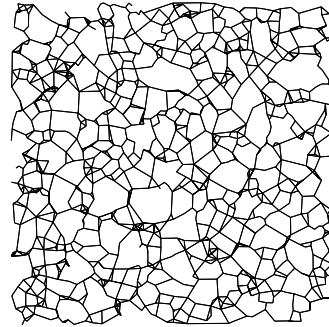
(c) NNTC ($k = 6$)



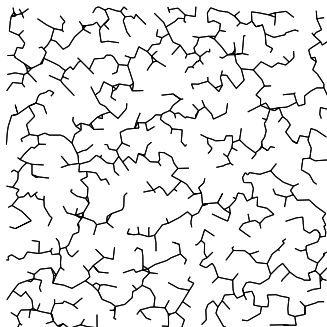
(d) Max Degree ($k = 6$)



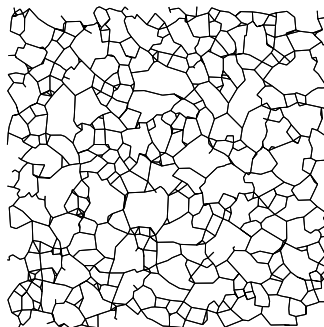
(e) CBTC, ohne Pairwise-Edge-Removal ($\alpha = 2\pi/3$)



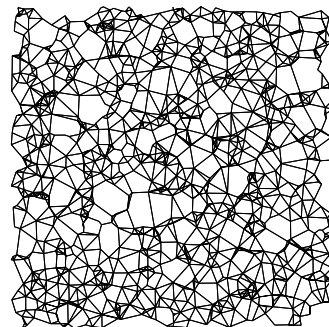
(f) CBTC ($\alpha = 2\pi/3$)



(g) MST



(h) RNG



(i) GG

Abbildung 6.1: Visualisierung von Beispieltopologien für unterschiedliche Formungsstrategien

Abbildung 6.2 : Zuordnung von Sendereichweiten in Max-Degree-Topologien

```

Parameter : Menge von Stationen  $S$ 
              Distanzfunktion  $d : S \times S \rightarrow \mathbb{R}$ 
              Nachbarzahl  $k \in \mathbb{N}$ 
              Maximale Sendereichweite  $r_{\max} \in \mathbb{R}$ 

forall  $v \in S$  do
   $r_v \leftarrow 0$ 
  repeat
     $I \leftarrow \{v \in S : r_v < r_{\max} \wedge |\{w \in S \setminus \{v\} : d_{vw} \leq \min\{r_v, r_w\}\}| < k\}$ 
    if  $I \neq \emptyset$  then
       $r_{\min} \leftarrow \min\{r_v : v \in I\}$ 
       $v \leftarrow$  zufälliges Element aus  $\{v \in I : r_v = r_{\min}\}$ 
       $r_v \leftarrow r_{\max}$ 
      forall  $w \in S$  do
        if  $r_{\min} < d_{vw} < r_v$  then
           $r_v \leftarrow d_{vw}$ 
    until  $I = \emptyset$ 
  forall  $v \in S : r_v = r_{\max}$  do
     $r_v \leftarrow \max\{d_{vw} : w \in S \wedge d_{vw} \leq r_w\}$ 

```

Stationen, die weniger als k bidirektionale Links haben und noch nicht das gesamte Netz mit ihrer Sendereichweite abdecken, diejenige mit der geringsten Sendereichweite ausgewählt. Ihre Sendereichweite wird so erhöht, dass sie genau eine Station mehr abdeckt (wie oben erläutert, kann davon ausgegangen werden, dass die Distanzen eindeutig sind). Diese Schleife wird so lange durchlaufen, bis alle Stationen k Links haben oder die maximale Sendereichweite erreicht ist. Abschließend wird die Sendereichweite jeder Station, für die keine k Links etabliert werden konnten, auf die Distanz zu dem am weitesten entfernten Nachbarn zurückgesetzt, zu dem ein bidirektionaler Link besteht.

Abbildung 6.1(d) zeigt die aus den Beispielpositionen resultierende Topologie für $k = 6$. Dort ist der in Abschnitt 4.2.1.1 beschriebene Nachteil zu erkennen, dass Links sich über sehr große Distanzen erstrecken können. Um dem etwas entgegenzuwirken, sind die Max-Degree-Topologien in den folgenden Auswertungen unter Berücksichtigung einer maximalen Sendereichweite von 200m erzeugt worden. Diese Entfernung ist im Verhältnis zur Stationsdichte immer noch recht groß, absolut betrachtet aber eine realistische Sendereichweite. Durch eine weitere Verringerung der maximalen Sendereichweite würde der beschriebene Nachteil der Max-Degree-Topologien weiter abgeschwächt. Es ist jedoch durchaus beabsichtigt, die negativen Auswirkungen der Topologeeigenschaften aufzuzeigen, die bei einer im Verhältnis zur Sendereichweite hohen Stationsdichte auftreten. Ohne die Beschränkung der Sendereichweiten auf 200m wären diese in den Auswertungen noch deutlich stärker ausgeprägt.

CBTC-Topologien Diese Topologien werden gemäß der Cone-Based-Strategie erzeugt (siehe Abschnitt 4.2.2.2). Diese ordnet jeder Station eine minimale Sendereichweite so zu, dass in der resultierenden Topologie jede Station einen Link in jedem Sektor mit einem Winkel von α hat. Danach werden die in Abschnitt 4.2.2.2 beschriebenen Optimierungen aus [WLBW01] angewendet.

Diese Optimierungsschritte sind nötig, weil die Topologien sonst sehr nachteilige Eigenschaften haben: Werden weder Asymmetric-Edge-Removal noch Pairwise-Edge-Removal durchgeführt, bauen die Stationen am Rand der Simulationsfläche so viele Links auf, dass die Topologie für das Beispielszenario mit 1000 Stationen nicht darstellbar ist (der durchschnittliche Knotengrad beträgt hier über 250!). Abbildung 6.1(e) zeigt die Topologie nach einem Asymmetric-Edge-Removal, es sind also all jene Links entfernt worden, die nicht von beiden beteiligten Stationen für eine möglichst große Abdeckung ihres Umkreises benötigt werden. Dadurch entfällt eine große Zahl von Links, es ist aber auch zu erkennen, dass die Stationen am Rand der Simulationsfläche untereinander Links über sehr große Distanzen aufrecht erhalten. Abbildung 6.1(f) zeigt die Topologie schließlich nach einem anschließenden Pairwise-Edge-Removal, der gut in der Lage ist, diese Links zu identifizieren und zu entfernen.

Die Durchführung der genannten Optimierungsschritte hat zur Folge, dass die aus unterschiedlichen Parameterbelegungen für α resultierenden Topologien sich kaum unterscheiden. In den folgenden Auswertungen wird daher nur der Fall $\alpha = 3$ betrachtet, also der Mindestwert, für den die Verbundenheit der Topologien garantiert ist.

MST-Topologien Eine MST-Topologie ergibt sich aus dem minimalen Spannbaum des Distanzgraphen (siehe Abschnitt 4.2.1.3): Die Sendereichweite einer Station ist hier die Distanz zum entferntesten Nachbarn im minimalen Spannbaum.

Es sollte betont werden, dass die hier untersuchten MST-Topologien im Allgemeinen keine minimalen Spannbäume sind, sondern dass sie vereinzelt Zyklen enthalten. Dies ist auch in der in Abbildung 6.1(g) dargestellten Beispieltopologie erkennbar. Dennoch gilt, dass die Sendereichweiten unter Gewährleistung der Verbundenheit des Netzes minimal sind.

RNG-Topologien Eine RNG-Topologie wird aus dem Relative-Neighbourhood-Graphen des Distanzgraphen gebildet (siehe Abschnitt 4.2.1.4). Abbildung 6.1(h) stellt eine solche Topologie beispielhaft dar.

GG-Topologien Diese Topologien werden schließlich aus den Gabriel-Graphen der Distanzgraphen erzeugt (siehe Abschnitt 4.2.1.5). Eine Beispieltopologie ist in Abbildung 6.1(i) dargestellt.

6.3 Verwendete Metriken

6.3.1 Konnektivität

Eine Fülle von Arbeiten beschäftigt sich mit der Wahrscheinlichkeit für die Konnektivität von Topologien (z.B. [GK98], [SBV01], [XK04], [Bet02], [YS03]). In diesen Arbeiten wird dazu die Wahrscheinlichkeit betrachtet, dass alle Knoten zur selben Zusammenhangskomponente gehören. Als Boole'sche Eigenschaft ist dieses Kriterium jedoch nicht in der Lage, den Grad der Konnektivität einer Topologien zu erfassen, die aus mehreren Zusammenhangskomponenten besteht: Beispielsweise können alle Stationen bis auf eine einzige, die außerhalb der Sendebereiche anderer Stationen liegt, zu einer großen Zusammenhangskomponente gehören. Eine solche Topologie ist gemäß dieses Konnektivitätskriteriums nicht verbunden, aber trotzdem ist die Wahrscheinlichkeit, dass zwei zufällig ausgewählte Stationen miteinander kommunizieren können, bei höherer Stationszahl nahezu 1.

Deshalb dient als Maß für die Konnektivität einer Topologie in der vorliegenden Arbeit die Wahrscheinlichkeit, dass zwei verschiedene, zufällig gewählte Stationen in derselben Zusammenhangskomponente liegen:

$$\text{Konnektivität} = \frac{|\{(v, w) \in S^2 : v \neq w \wedge |R_{vw}| > 0\}|}{|S| \cdot (|S| - 1)},$$

wobei R_{vw} die Menge möglicher Routen zwischen v und w bezeichnet.

Dieses Konnektivitätsmaß wird auch in [DTH02] verwendet.

6.3.2 Routenoverhead

Der Overhead, der durch unnötig lange Kommunikationspfade entsteht, wird häufig mit Hilfe des *Dehnungsfaktors* gemessen, der bereits in Abschnitt 2.4.2 eingeführt wurde. Er gibt die maximale *Dehnung* über alle Knotenpaare des Graphen an, wobei die Dehnung zwischen zwei Knoten definiert ist als der Quotient aus der Länge des kürzesten Pfads und der euklidischen Distanz zwischen den beiden Knoten. In diesem Zusammenhang wurden auch zwei Eigenschaften erläutert, die die Eignung dieser Metrik für die folgenden Auswertungen einschränken. Zunächst geht in den Dehnungsfaktor nur ein Quelle-Senke-Paar ein, nämlich das, welches mit dem größten Overhead verbunden ist, während für die folgenden Auswertungen ein Mittelwert über alle Quelle-Senke-Paare als Maß für den tatsächlich anfallenden Overhead geeigneter ist. Außerdem ist die Gesamtdistanz, die beim Transport eines Pakets zurückgelegt wird, durch die Summe der Sendereichweiten entlang der Route gegeben; der Overhead, der durch unnötig hohe Sendeleistungen verursacht wird, ist im Dehnungsfaktor daher nicht enthalten.

Die Quantifizierung des Overheads, der mit einer Netztopologie verbunden ist, lehnt sich in der vorliegenden Arbeit zwar an den Dehnungsbegriff an, die im Folgenden eingeführte Metrik wird jedoch als *Routenoverhead* bezeichnet, um Missverständnissen vorzubeugen. Sie

gibt die mittlere Dehnung über alle Knotenpaare an, wobei eine Kantenlänge aber nicht durch die euklidische Distanz zwischen den beiden Knoten gegeben ist, sondern durch die Sendeleistung des Knotens, von dem die Kante ausgeht. Der *Routenoverhead* ist also definiert durch

$$\text{Routenoverhead} = \left(\prod_{v \in S} \prod_{w \in S \setminus \{v\}} \sum_{u \in S \setminus \{w\}} \frac{|R_{vw} \cap R_u|}{|R_{vw}|} \cdot r_u / d_{vw} \right)^{1/(|S| \cdot (|S|-1))},$$

wobei r_u die Sendereichweite einer Station $u \in S$ bezeichnet und d_{vw} die euklidische Distanz zwischen zwei Stationen $v, w \in S$. Weiterhin enthält die Menge R_{vw} alle Pfade zwischen v und w , die gemäß einer vorgegebenen Routingmetrik minimale Länge haben, und R_u ist die Menge aller Pfade, die u enthalten. Die Verwendung alternativer Pfade einer Menge R_{vw} wird in dieser Metrik damit als gleich wahrscheinlich angenommen.

Der Routenoverhead bezeichnet also die über alle Quelle-Senke-Paare vorgenommene geometrische Mittelung der Quotienten aus der „Brutto-Distanz“ und der „Netto-Distanz“ zwischen Quelle und Senke: Die Brutto-Distanz, d.h. die Entfernung, über die ein Paket transportiert wird, ist der Erwartungswert für die Summe der Sendereichweiten aller Stationen entlang einer Route. Die Netto-Distanz ist die euklidische Distanz zwischen Quelle und Senke.

Da hier eine Mittelung von Werten vorgenommen wird, die Größenverhältnisse angeben, wäre ein arithmetisches Mittel nicht aussagekräftig. Alternativ zum geometrischen Mittel kann auch der Quotient aus den über alle Quelle-Senke-Paare gebildeten Erwartungswerten für Brutto- und Netto-Distanz betrachtet werden:

$$\text{Gewichteter Routenoverhead} = \frac{\sum_{v \in S} \sum_{w \in S \setminus \{v\}} \sum_{u \in S \setminus \{w\}} \frac{|R_{vw} \cap R_u|}{|R_{vw}|} \cdot r_u}{\sum_{v \in S} \sum_{w \in S \setminus \{v\}} d_{vw}}.$$

Dabei ist allerdings problematisch, dass größere Distanzen in diesem Wert auch ein stärkeres Gewicht haben, obwohl eine Gewichtung nur dann sinnvoll wäre, wenn diese der Aufteilung der Ressourcen unter den Strömen entspricht. In Zähler und Nenner werden nämlich die Distanzen aufsummiert, die zurückgelegt werden, wenn jeder Strom genau ein Paket von der Quelle zum Ziel transportiert, also unter Annahme strikter Fairness. Der Fairnessbegriff, der den folgenden Auswertungen zugrunde liegt, ist jedoch der der Max-Min-Fairness (siehe Abschnitt 2.6). Da also ohne die konkrete Berechnung eines TDMA-Schemas keine Aussagen über die Ressourcenverteilung auf die Ströme gemacht werden kann, wird der Routenoverhead hier durch das geometrische Mittel berechnet, in das jedes Quelle-Senke-Paar mit dem gleichen Gewicht eingeht.

Im Übrigen haben Vergleiche gezeigt, dass der gewichtete Overhead nur in solchen Szenarien eine bessere Übereinstimmung mit dem tatsächlichen Overhead zeigt, in denen kein Spatial-Reuse möglich ist, so dass sich Max-Min-Fairness nicht von strikter Fairness unterscheidet. Es handelt sich also um die weniger interessanten Szenarien, für die diese Metrik im Folgenden auch gar nicht verwendet wird. In allen anderen Szenarien ist die Übereinstimmung des als geometrisches Mittel definierten Routenoverheads dagegen sehr viel genauer.

In den folgenden Auswertungen werden zwei Routingmetriken in Betracht gezogen: Zum einen wird der Optimalfall untersucht, dass die Routingmetrik die Summe der Sendereichweiten entlang des Pfads zwischen Quelle und Senke minimiert, zum anderen der in Anbetracht aktueller Routingprotokolle realistischere Fall, dass die Anzahl der Übertragungen entlang des Pfads minimiert wird.

6.3.3 Inhomogenität von Sendeleistungszuweisungen

Sind die Sendeleistungen von Stationen, die sich in räumlicher Nähe zueinander befinden, stark verschieden, so kann der Medienzugang unter Umständen weniger effektiv durch das Data-Link-Protokoll organisiert werden. Beispielsweise wird das in Abschnitt 2.2.1 beschriebene Hidden-Node-Problem durch inhomogene Sendeleistungen verstärkt, weil es wahrscheinlicher ist, dass ein potenzieller Störer nicht in der Lage ist, das Signal des Senders wahrzunehmen oder ein RTS oder CTS zu empfangen (dies wurde bereits in Abschnitt 3.3 angesprochen und in Abbildung 3.2 skizziert).

Zur Messung dieser Inhomogenität der Sendeleistungszuweisungen kann im einfachsten Fall die Anzahl unidirektionaler Links dienen. Wie in Abschnitt 2.4.1 erläutert, ist ein unidirektionaler Link eine Kante (v, w) im Topologiegraphen, deren zugehörige Kante in Rückrichtung (w, v) nicht im Topologiegraphen enthalten ist. Da höhere Sendereichweiten im Allgemeinen mit einer höheren Anzahl sowohl unidirektionaler als auch bidirektionaler Links einhergehen, kann es auch sinnvoll sein, die Relation dieser Zahlen zu betrachten. In dieser Arbeit wird dazu die *Asymmetrie* einer Topologie definiert durch

$$\text{Asymmetrie} = \frac{|\{(v, w) \in S^2 : v \neq w \wedge r_v \geq d_{vw}\}|}{|\{(v, w) \in S^2 : v \neq w \wedge \min\{r_v, r_w\} \geq d_{vw}\}|}.$$

Der Zähler dieses Bruchs ist die Anzahl aller gerichteten Kanten im Topologiegraphen, also auch jener Kanten, die unidirektionale Links sind. Im Nenner werden hingegen nur die gerichteten Kanten gezählt, die zu bidirektionalen Links gehören. Eine Kante (v, w) , die im Topologiegraphen genau dann enthalten ist, wenn sich w in Sendereichweite von v befindet, wird im Zähler also in jedem Fall mitgezählt, im Nenner aber nur dann, wenn v sich ebenfalls in Sendereichweite von w befindet, wenn die Kante also zu einem bidirektionalen Link gehört. Bidirektionale Links werden daher doppelt gezählt: Sowohl (v, w) als auch (w, v) gehen als jeweils eine Kante in die Metrik ein.

Die Asymmetrie ist so definiert, dass sie das Verhältnis zwischen dem Erwartungswert der Anzahl eingehender Links einer Station und dem Erwartungswert der Anzahl bidirektionaler Links einer Station angibt. Beträgt die durchschnittliche Anzahl bidirektionaler Links einer Station also beispielsweise 10, so folgt aus einem Asymmetrie-Wert von 1,1, dass an jeder Station im Schnitt 11 Kanten eingehen. In diesem Fall wäre jede Station im Schnitt also von einem eingehenden unidirektionalen Link betroffen.

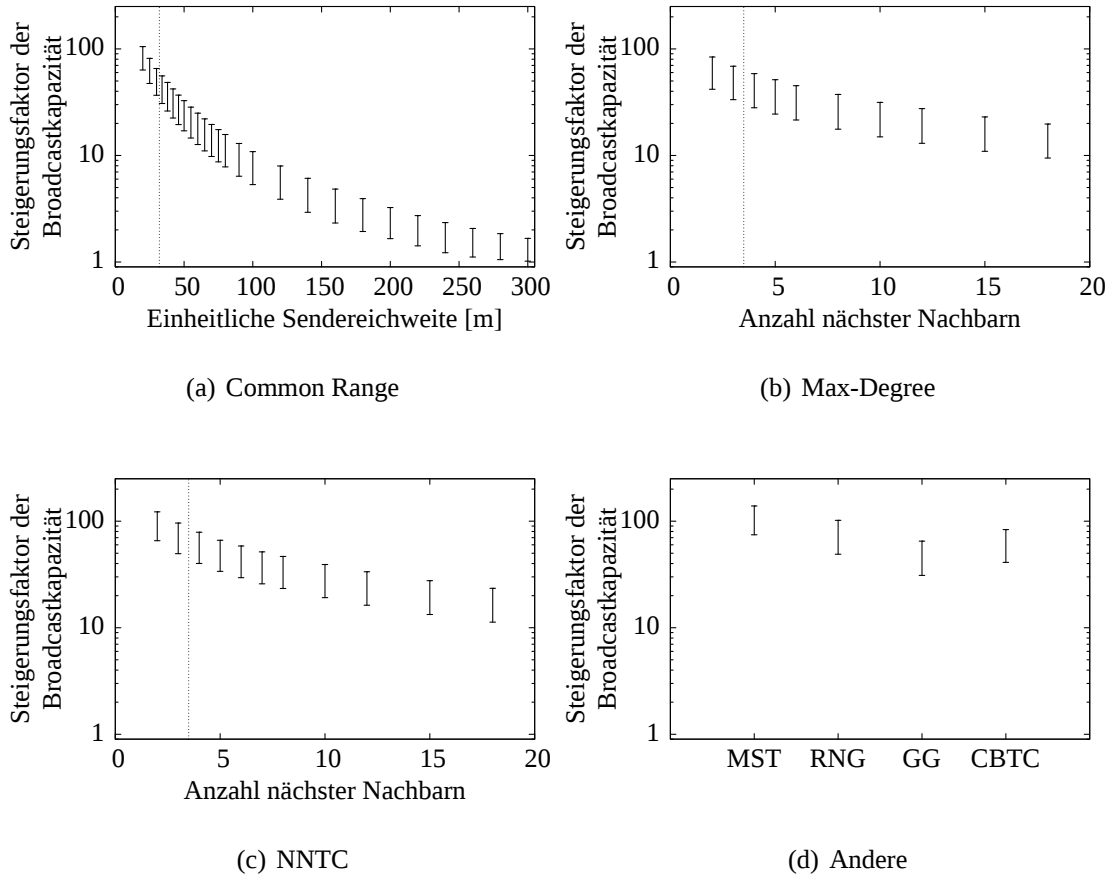


Abbildung 6.3: Broadcast-Kapazität mit verschiedenen Strategien zur Sendeleistungsregelung (obere Intervallgrenze $\delta = 0$, untere Intervallgrenze $\delta = 1$)

6.4 Broadcast-Kapazität

Im folgenden Abschnitt wird die Broadcast-Kapazität mit Hilfe des in Abschnitt 5.2 vorgestellten Verfahrens gemessen.

Abbildung 6.3 zeigt, um welchen Faktor sich die Broadcast-Kapazität gegenüber einem vollständig verbundenen Netz für verschiedene Strategien zur Sendeleistungsregelung im Mittel erhöht. In den 100 betrachteten Szenarien sind 1000 Stationen zufällig und gleichverteilt auf einer quadratischen Fläche mit einer Diagonalen von 1000m positioniert. Die Intervallgrenzen korrespondieren mit zwei verschiedenen Belegungen des Parameters δ : Hierbei stellt $\delta = 0$ einen sehr optimistischen Fall dar, während es sich bei $\delta = 1$ um eine eher vorsichtige Einschätzung handelt. (Ersterer korrespondiert also mit der Obergrenze, letzterer mit der Untergrenze der dargestellten Intervalle.) Die 95%-Konfidenzintervalle der dargestellten Mittelwerte sind für beide δ vernachlässigbar gering. Für einige der Strategien führt eine unge-

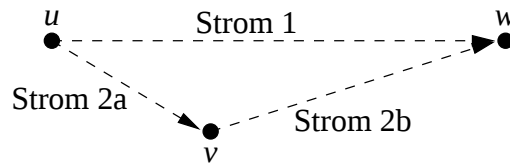


Abbildung 6.4: Zweiteilung eines langen Datenstroms

schickte Parameterwahl zwar zu Topologien mit einer sehr hohen Broadcast-Kapazität, die aber eine geringe Konnektivität aufweisen. In den Grafiken kennzeichnen vertikale Linien deshalb die Untergrenze des Bereichs, in dem Topologien eine Konnektivität von mindestens 95% aufweisen.

Es lässt sich erkennen, dass alle Strategien eine deutliche Steigerung des Spatial-Reuse und damit der Broadcast-Kapazität gegenüber einem vollständig verbundenen Netz ermöglichen. Mit höheren Sendereichweiten nimmt die Kapazität rapide ab. Die höchsten Werte ergeben sich für die MST-Topologien, was dadurch zu erklären ist, dass sie minimale Sendereichweiten bei maximaler Konnektivität garantieren.

Werden sämtliche Daten ohne die explizite Verwaltung von Kommunikationspfaden geflutet, ist es also vorteilhaft, möglichst geringe Sendereichweiten einzustellen. Es ist lediglich zu berücksichtigen, dass das Netz verbunden bleibt. Dies ist auch zu erwarten gewesen, da die Gesamtzahl der Übertragungen immer gleich ist und sich lediglich der Spatial-Reuse verändert. Komplizierter sind die Auswirkungen der eingestellten Sendereichweiten auf die Transportkapazität, da eine Verringerung der Reichweiten dort aufgrund verlängerter Kommunikationspfade auch eine höhere Anzahl von Einzelübertragungen impliziert. Dies wird im folgenden Abschnitt ausgeführt.

6.5 Transportkapazität

Im folgenden Abschnitt wird der Zusammenhang zwischen den Sendereichweiten der Geräte eines drahtlosen Netzes und dessen Transportkapazität untersucht. Dazu werden in Abschnitt 6.5.1 zunächst Netze betrachtet, in denen allen Stationen die gleiche Sendereichweite zugeordnet ist. Abschnitt 6.5.2 widmet sich dem Vergleich von NNTC- und Max-Degree-Topologien, und in Abschnitt 6.5.3 wird schließlich kurz auf andere Topologieklassen eingegangen.

Dazu wird der erzielbare Gesamtdurchsatz mit Hilfe des in Abschnitt 5.3 eingeführten Zwei-Phasen-Färbalgorithmus approximiert und in der Einheit *Paket-Meter pro Zeitslot* angegeben. In diesem Wert wird der Durchsatz jedes Stroms also mit der Distanz zwischen Quelle und Senke gewichtet, da das Überbrücken größerer Entfernung mit einem stärkeren Ressourcenverbrauch verbunden ist. Dieser *gewichtete Gesamtdurchsatz* gibt daher an, welche Datenmenge in welcher Zeit über welche Distanz transportiert wird.

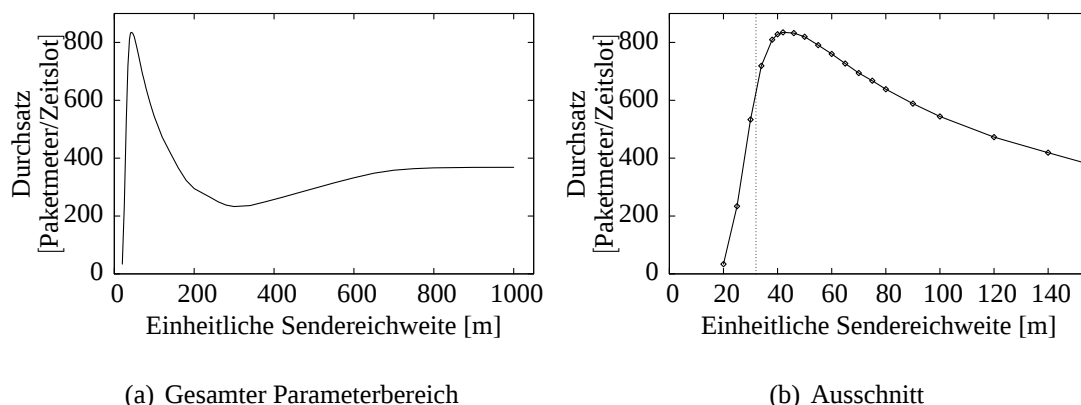


Abbildung 6.5: Transportkapazität in Common-Range-Topologien mit 1000 Stationen

Wird dies wie bei der einfachen Addition der Durchsatzwerte zu einem ungewichteten Gesamtdurchsatz nicht berücksichtigt, so sind die Ergebnisse noch sehr viel stärker vom Verkehrsmuster abhängig, als es gemäß der in Abschnitt 6.1.3 geführten Argumentation der Fall ist. So wird der Anteil eines Datenstroms, der von einer Quelle u über eine Zwischenstation v zu einer Senke w führt, am ungewichteten Gesamtdurchsatz mindestens verdoppelt, wenn dieser Strom an v zweigeteilt wird, d.h. durch zwei Ströme ersetzt wird, die Daten von u nach v und von v nach w transportieren (siehe Abbildung 6.4). Tatsächlich steigt der ungewichtete Gesamtdurchsatz noch stärker, wenn einer der resultierenden Ströme einen höheren Durchsatz erzielen kann, weil sein Pfad nicht den Flaschenhals enthält, der die Bandbreite des ursprünglichen Stroms begrenzt hat. Im Gegensatz zur zuvor beschriebenen Verdopplung ist dies keine willkürliche Erhöhung, sondern dadurch zu erklären, dass es mit kürzeren Quelle-Senke-Distanzen besser möglich ist, vorhandene Ressourcen vollständig zu nutzen.

Daher wirkt sich eine solche Zweiteilung eines Stroms im Allgemeinen auch bei der durchgeführten Gewichtung der einzelnen Durchsätze mit den Quelle-Senke-Distanzen auf den gewichteten Gesamtdurchsatz aus, jedoch weniger willkürlich als beim ungewichteten Gesamtdurchsatz. Neben dem bereits genannten Grund, dass einer der beiden kürzeren Ströme einen höheren Durchsatz erreichen kann, kann sich auch die Summe der Quelle-Senke-Distanzen gemäß der Dreiecksungleichung vergrößern.

6.5.1 Common-Range-Topologien

In Abbildung 6.5 ist der mittlere Gesamtdurchsatz von 250 Strömen in den bereits im letzten Abschnitt untersuchten Common-Range-Topologien aufgetragen, in denen 1000 Stationen gleichverteilt auf einer quadratischen Fläche mit einer Diagonalen von 1000m platziert sind. Die Schwankungen der Resultate sind so gering, dass die 95%-Konfidenzintervalle für den Erwartungswert vernachlässigbar klein sind und deshalb nicht aufgetragen sind.

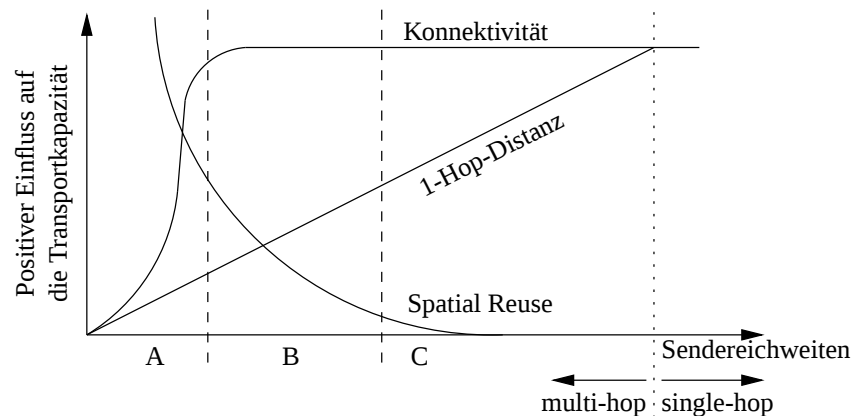


Abbildung 6.6: Einfluss verschiedener Größen auf die Transportkapazität (Skizze)

Abbildung 6.5(a) zeigt den Verlauf der Transportkapazität über den gesamten Bereich der Sendereichweiten, die einen Einfluss auf die resultierende Netztopologie haben: Reichweiten über 1000m sind zwar möglich, jedoch ergibt sich für jede einheitliche Sendereichweite von mindestens 1000m immer eine vollständig verbundene Netztopologie. Dieser Verlauf der Transportkapazität lässt drei Bereiche erkennen.

Im Bereich sehr geringer Sendereichweiten (bis ca. 40m) wirkt sich eine Steigerung der Sendereichweiten positiv auf die Netzkapazität aus, da zunächst wegen der Erhöhung der Konnektivität mehr Ströme überhaupt zustande kommen können. Außerdem haben auch die Längen der Pfade zwischen Quellen und Senken in diesem Bereich einen gewissen Einfluss auf die Transportkapazität; dies wird weiter unten erläutert.

Ebenfalls im Bereich sehr hoher Sendereichweiten (ab ca. 300m) steigt die Kapazität mit steigenden Sendereichweiten, weil durch die Verkürzung der Routen weniger einzelne Übertragungen erforderlich sind, um ein Paket von einer Quelle zu einer Senke zu transportieren.

Zwischen diesen beiden Bereichen liegt ein dritter Bereich, in dem sich eine Erhöhung der Sendereichweiten negativ auf die Transportkapazität auswirkt. Offensichtlich überwiegt hier der Einfluss des Spatial-Reuse: Die Anzahl der Übertragungen, die gleichzeitig stattfinden können, wird durch die Erhöhung der Sendereichweiten so stark reduziert, dass dies auch durch die verkürzten Routen und die damit verbundene geringere Anzahl von Übertragungen nicht ausgeglichen wird. Im Bereich höherer Sendereichweiten ändert sich dieses Verhältnis, weil schon für Sendereichweiten, die deutlich kleiner sind als die maximale Distanz zwischen zwei Stationen, praktisch kein Spatial-Reuse mehr möglich ist.

Dieser beschriebene Einfluss der verschiedenen Größen auf die Transportkapazität ist in Abbildung 6.6 skizziert, wobei die Grenzen zwischen den Bereichen als fließend anzusehen sind.

Der Übergang vom ersten, von der Konnektivität dominierten zum zweiten, vom Spatial-Reuse dominierten Bereich wird durch die Stationsdichte bestimmt. Ist also die Größe und die Form einer Fläche fest vorgegeben, so wird eine ausreichende Konnektivität mit umso geringeren Sendereichweiten erzielt, je mehr Stationen auf dieser Fläche platziert sind, und

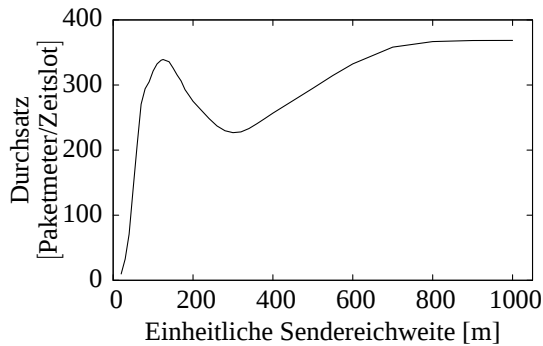


Abbildung 6.7: Transportkapazität in Szenarien mit 100 Stationen

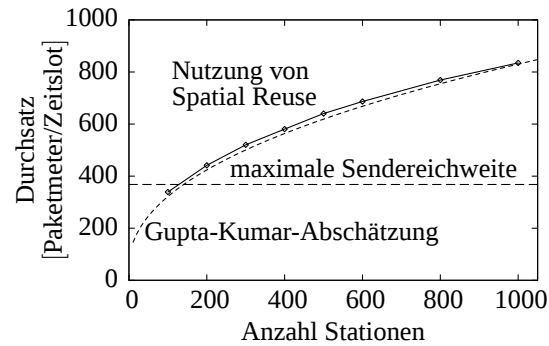


Abbildung 6.8: Maximal erzielbarer Durchsatz in Common-Range-Topologien in Abhängigkeit von der Stationszahl

prinzipiell kann die für ausreichende Konnektivität benötigte Sendereichweite durch das Hinzufügen von immer mehr Stationen beliebig nahe an 0 gebracht werden.

Der Übergang vom zweiten, vom Spatial-Reuse dominierten in den dritten, von den Routenlängen dominierten Bereich hängt hingegen von der Größe (und der Form) der Fläche an sich ab, denn unabhängig von der konkreten Stationszahl blockiert eine laufende Übertragung ab einer gewissen Sendereichweite zumindest einen großen Teil des Netzes.

Aus diesen Feststellungen folgt, dass der Spatial-Reuse umso besser genutzt werden kann, je mehr Stationen das Netz umfasst. Bei einer geringen Zahl von Stationen spielt der Spatial-Reuse hingegen schon keine Rolle mehr, wenn eine ausreichende Konnektivität erreicht ist. Dies wird durch Abbildung 6.7 verdeutlicht, in der der mittlere Durchsatz unter den gleichen Bedingungen aufgetragen ist, die Abbildung 6.5(a) zugrunde liegen, mit Ausnahme der geringeren Stationszahl von 100. Hier ist zwar noch ein Bereich sichtbar, in dem der Einfluss des Spatial-Reuse dominiert, jedoch sind die durch optimale Ausnutzung des Spatial-Reuse erzielbaren Verbesserungen vergleichsweise gering und bringen keinen Vorteil gegenüber einem vollständig verbundenen Netz.

Abbildung 6.8 zeigt den mittleren Gesamtdurchsatz in Abhängigkeit der Stationszahl, der sich mit einheitlichen Sendereichweiten und bei Nutzung von Spatial-Reuse erzielen lässt. Die Sendereichweite wurde also im Übergang zwischen den in Abbildung 6.6 mit A und B bezeichneten Bereichen (in denen die Konnektivität bzw. der Spatial-Reuse dominieren) so gewählt, dass der resultierende Gesamtdurchsatz möglichst hoch ist. Zum Vergleich ist auch der mit maximaler Sendereichweite erzielbare Durchsatz aufgetragen, der unabhängig von der Stationszahl ist, da hier in jedem Zeitslot genau ein Paket von seiner Quelle zu seiner Senke transportiert wird. Der Durchsatz bestimmt sich in diesem Fall daher nur durch den mittleren Abstand zwischen Quelle und Senke (hier ca. 368m). Außerdem enthält die Abbildung die Abschätzung der Transportkapazität nach [GK00] (siehe Abschnitt 4.1.1)

$$n \cdot (cW / \sqrt{n \log n}) \cdot 368\text{m},$$

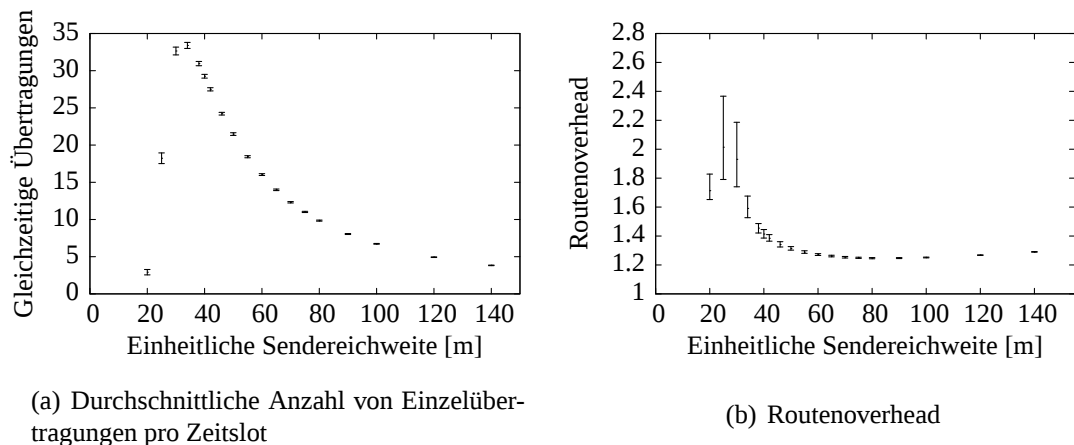


Abbildung 6.9: Spatial-Reuse und Routenoverhead in Common-Range-Topologien mit 1000 Stationen

für $c = 0,1875$ und $W = 1\text{Paket}/\text{Zeitslot}$. Obwohl dieser Abschätzung andere Annahmen zugrunde liegen, ist eine relativ gute Übereinstimmung des Wachstums an Gesamtkapazität mit steigender Stationszahl zu erkennen.

Diese Tatsachen erklären im Übrigen, dass sich immer wieder Arbeiten finden, die anhand von Simulationsergebnissen begründen, dass der höchste Durchsatz grundsätzlich durch größtmögliche Sendeleistungen erzielt wird, wie z.B. die in Abschnitt 4.2.3.2 vorgestellte Arbeit [PS02a], die sich auf Simulationsszenarien mit höchstens 100 Stationen bezieht. Der in Abbildung 6.7 erkennbare Einfluss des Spatial-Reuse für Sendereichweiten zwischen ca. 150m und 300m ist anhand simulativer Auswertungen kaum festzustellen, weil die modellierten Protokolle das Potenzial des Spatial-Reuse nicht gut nutzen können und weil die Ergebnisse solcher Simulationen generell auch sehr viel stärkeren Schwankungen unterworfen sind.

Wie eingangs bereits erwähnt, haben auch die Pfade zwischen Quelle-Senke-Paaren unter bestimmten Umständen einen nicht unerheblichen Einfluss auf die Transportkapazität des Netzes. Dies ist für den Bereich von Sendereichweiten der Fall, in dem die meisten Quellen-Senke-Paare zwar verbunden sind, eine ausreichende Konnektivität also gegeben ist, in dem aber noch so wenig Links vorhanden sind, dass eine Multihop-Verbindung nur über große Umwege möglich ist. Dieser Bereich ist in Abbildung 6.5(b) gut zu erkennen: Obwohl die Konnektivität der Topologien schon für einheitliche Sendeleistungen um 40m nahezu 1 ist, bewirkt eine Erhöhung der Sendereichweiten bis 50m keine nennenswerte Einbuße der Netzkapazität, die in diesem Parameterbereich um ca. 800 Paketmeter pro Zeitslot liegt. Dabei zeigt Abbildung 6.9(a), in der die durchschnittliche Anzahl durchgeführter Übertragungen pro Zeitslot dargestellt ist, dass sich eine Änderung der Sendereichweiten auch in diesem Bereich deutlich auf den Spatial-Reuse auswirkt: Während bei einer Sendereichweite von 50m im Mittel 21,5 Übertragungen gleichzeitig stattfinden, pro Zeitslot also $21,5 \cdot 50\text{m} = 1075\text{m}$ zurückgelegt werden, sind es bei 38m im Mittel 31 gleichzeitige Übertragungen, so dass pro

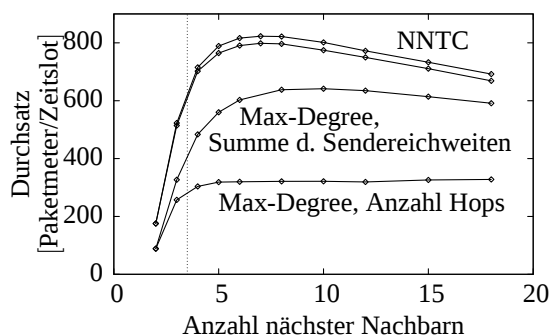


Abbildung 6.10: Transportkapazität von NNTC- und Max-Degree-Topologien mit 1000 Stationen

Zeitslot $31 \cdot 38\text{m} = 1178\text{m}$ zurückgelegt werden. Dies entspricht einer Steigerung von fast 10%. Die Routenoverhead-Metrik, die in Abschnitt 6.3.2 erläutert wurde und deren 5- und 95-Perzentile für die verschiedenen Topologien in Abbildung 6.9(b) aufgetragen sind, zeigt aber, wie stark auch das Verhältnis zwischen Brutto- und Netto-Distanzen hier schwankt. Bei einer einheitlichen Sendereichweite von 50m beträgt der Routenoverhead im Mittel ca. 1,31, bei 38m schon ca. 1,45. Durch diese Steigerung um etwas mehr als 10% wird die Erhöhung des Spatial-Reuse also ungefähr wieder ausgeglichen.

6.5.2 NNTC- und Max-Degree-Topologien

Für die im vorangegangenen Abschnitt besprochenen Common-Range-Topologien spielt es keine Rolle, ob kürzeste Routen gemäß der Anzahl der Übertragungen oder gemäß der Summe der Sendereichweiten entlang der Route gewählt werden. Da die Sendereichweiten aller Stationen gleich sind, unterscheiden sich beide Metriken nur um genau die Sendereichweite als konstanten Faktor. Werden den Stationen aber individuelle Sendeleistungen zugewiesen, macht es durchaus einen Unterschied, welche Routingmetrik verwendet wird. Die Minimierung der Summe der Sendereichweiten verspricht dabei einen höheren Spatial-Reuse.

Abbildung 6.10 zeigt den erreichbaren Durchsatz für NNTC- und Max-Degree-Topologien mit beiden Routingmetriken. Für beide Topologieklassen liegt die Konnektivität ab einer Nachbarzahl von $k = 4$ im Mittel über 95%. Der Verlauf des erzielbaren Durchsatzes in Abhängigkeit von der Nachbarzahl ist qualitativ vergleichbar zu dem in Abbildung 6.5(b) dargestellten Durchsatz in Common-Range-Topologien. Die im vorangegangenen Abschnitt ausführlich dargelegten Begründungen für diesen Verlauf gelten hier genauso, da die Sendereichweiten mit der vorgegebenen Nachbarzahl monoton steigen.

Der mit NNTC maximal erreichbare Durchsatz unterscheidet sich unabhängig von der verwendeten Routingmetrik nicht wesentlich von dem, der mit einer einheitlichen Sendereichweite im Bestfall zu erreichen ist. Wird die Länge eines Pfads durch die Anzahl der Übertragungen entlang des Pfads gemessen, so ist der Gesamtdurchsatz nur geringfügig niedriger,

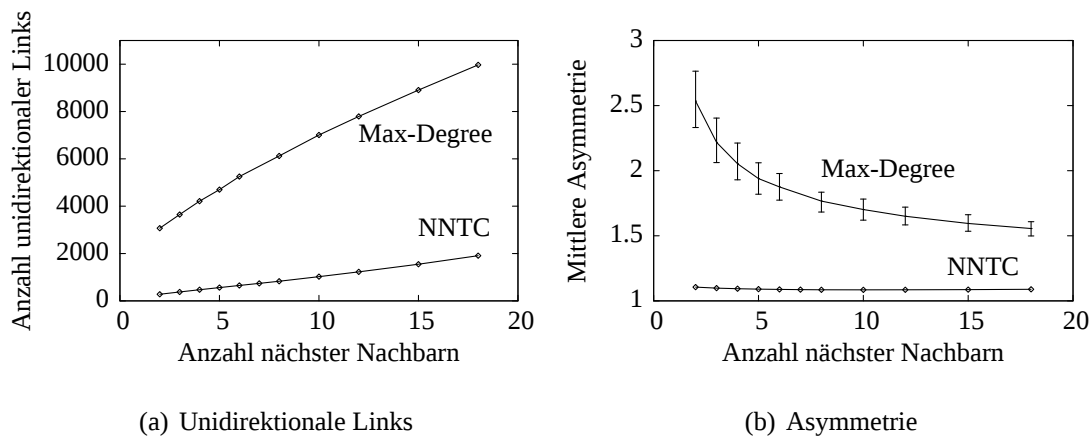


Abbildung 6.11: Inhomogenität der Sendereichweiten in NNTC- und Max-Degree-Topologien mit 1000 Stationen

als wenn sie durch die Summe der Sendereichweiten entlang des Pfads gemessen wird. Für Max-Degree-Topologien liegt der Durchsatz in beiden Fällen auf deutlich geringerem Niveau, und der Einfluss der gewählten Routingmetrik ist sehr stark ausgeprägt.

Dies ist durch den in Abschnitt 4.2.1.1 beschriebenen Nachteil der Max-Degree-Topologien zu begründen: Während die Sendereichweiten in Umgebungen mit hoher Stationsdichte eher gering bleiben, weil das vorrangige Ziel die Einhaltung einer Obergrenze für die Nachbarzahl ist, müssen andere Stationen Links über unnötig weite Distanzen aufbauen, um die vorgegebene Nachbarzahl erreichen zu können. Führen solche Stationen Übertragungen durch, blockieren sie eine große Anzahl anderer Übertragungen. Falls die Sendereichweiten der Stationen also in die Routingmetrik eingehen, werden besonders ungünstige Zwischenstationen vermieden, nicht jedoch, wenn die Routingmetrik nur die Anzahl der Zwischenstationen berücksichtigt.

Diese stark ausgeprägte Inhomogenität der Sendeleistungszuweisungen wird im Folgenden genauer betrachtet. Abbildung 6.11(a) zeigt die mittlere Zahl unidirektionaler Links in den betrachteten Szenarien (die 95%-Konfidenzintervalle für den Mittelwert sind so klein, dass sie nicht dargestellt werden). Dort ist zu sehen, dass die Zahl unidirektionaler Links für beide Topologieklassen mit steigender Nachbarzahl zunimmt, dass Max-Degree-Topologien für eine gegebene Nachbarzahl aber um ein Vielfaches mehr unidirektionale Links enthalten als NNTC-Topologien.

In Abbildung 6.11(b) sind die geometrischen Mittel der Asymmetrie-Werte gemäß der in Abschnitt 6.3.3 vorgestellten Metrik beider Topologieklassen dargestellt. Da diese Werte für Max-Degree-Topologien starken Schwankungen unterworfen sind, sind auch die 5- und 95-Perzentile dargestellt, die für NNTC-Topologien nur geringfügig vom Mittelwert abweichen. Während diese für NNTC-Topologien fast unabhängig von der Nachbarzahl k zumeist unter 1,1 liegen, zeigen die Asymmetrie-Werte für Max-Degree-Topologien eine deutliche Abhän-

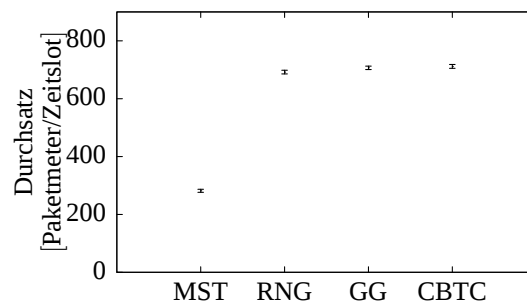


Abbildung 6.12: Transportkapazität in anderen Topologien

gigkeit von k . Mit steigendem k nimmt die Asymmetrie ab, weil die Wahrscheinlichkeit dafür, dass die maximale Distanz innerhalb einer Gruppe von $k + 1$ Stationen geringer ist als die minimale Distanz zu einer anderen Station (wie in Abbildung 4.2 veranschaulicht), mit steigendem k abnimmt. Dennoch liegt die Asymmetrie von Max-Degree-Topologien für $k = 18$ im Mittel immer noch bei über 1,5.

6.5.3 Andere Topologien

Abbildung 6.12 zeigt den erzielbaren Durchsatz in Topologien, die gemäß anderer Strategien geformt wurden. In MST-Topologien ist der Durchsatz vergleichsweise gering, was sich durch den extrem hohen Routenoverhead begründen lässt, der im Mittel bei über 3 liegt; Pakete werden im Mittel also physisch über Entfernungen transportiert, die mehr als drei mal so groß sind wie die Distanzen zwischen Sendern und Empfängern. Zum Vergleich sei daran erinnert, dass dieser Wert in Common-Range-Topologien mit gerade ausreichender Konnektivität 1,6 beträgt und mit etwas höheren Sendereichweiten schnell Werte um 1,3 erreicht (siehe Abbildung 6.9(b)).

Die anderen Topologien weisen eine vergleichbare Transportkapazität auf, die zwar nicht wesentlich geringer, aber auch nicht höher ist als diejenige, die in Common-Range- oder NNTC-Topologien mit einer günstigen Parameterwahl zu beobachten ist. Wie in Kapitel 4 beschrieben wurde, sind diese Topologieklassen aus praktischen Gründen als Basis eines verteilten Verfahrens zur Sendeleistungsregelung in Geräten ohne spezielle Hardware jedoch nicht geeignet.

6.6 Fazit

In den vorangegangenen Abschnitten wurden Topologien bewertet, die nach unterschiedlichen Strategien geformt wurden. Diese Strategien benötigen (mit Ausnahme von CBTC) keine geometrische Information (siehe Abschnitt 4.2).

Es hat sich gezeigt, dass auf der relativ abstrakten Ebene, auf der die Untersuchungen durchgeführt wurden, ein deutlicher Kapazitätsgewinn durch die Regelung von Sendeleistungen erzielt werden kann. Die Broadcast-Kapazität lässt sich in den betrachteten Szenarien mit 1000 Stationen gegenüber einem vollständig verbundenen Netz um Faktoren der Größenordnung 100 steigern. Auch die Transportkapazität kann deutlich gesteigert werden (bei 1000 Stationen auf über das Doppelte), wobei Common-Range- und Nearest-Neighbours-Topologien bei geeigneter Parameterwahl den höchsten Durchsatz erlauben.

Max-Degree-Topologien lassen hingegen den bereits in Abschnitt 4.2.1.1 beschriebenen Nachteil auch quantitativ deutlich erkennen: Die Sendereichweiten der Stationen im Netz sind sehr viel inhomogener als in Nearest-Neighbours-Topologien. Dies hat zunächst praktische Nachteile für die Regelung des Medienzugangs. Außerdem bleibt weniger durch Spatial-Reuse nutzbares Potenzial, insbesondere wenn die verwendete Routingmetrik die Sendereichweiten der Stationen nicht berücksichtigt und damit nicht in der Lage ist, Stationen mit hohen Sendereichweiten nach Möglichkeit zu umgehen.

Es zeigt sich auch, dass verschiedene andere Topologieklassen, die eine maximale Konnektivität garantieren, in Bezug auf die Kapazität und damit verwandten Eigenschaften keine Vorteile gegenüber einfachen Strategien bieten, sofern deren Parameterbelegung geschickt gewählt ist.

7 Ein verteiltes Verfahren zur Sendeleistungsregelung

7.1 Motivation

Die im vorangegangenen Kapitel präsentierten Ergebnisse zeigen, dass die Kapazität großer und dichter drahtloser Netze deutlich gesteigert werden kann, indem die Sendeleistungen der beteiligten Geräte in geeigneter Weise geregelt werden. Die höchste Kapazität wiesen dabei zum einen die Netze auf, in denen die Stationen eine günstig gewählte einheitliche Sendereichweite hatten, und zum anderen die Netze, in denen die Sendereichweiten gemäß der Nearest-Neighbours-Strategie mit geeignetem Nachbarzahlparameter bestimmt wurden.

Die Nearest-Neighbours-Strategie hat zunächst den Vorteil, dass sie dezentral und nur mit Information über die lokale Umgebung eines Geräts umgesetzt werden kann, weil der Parameter für die Mindestnachbarzahl k dafür sorgt, dass die Sendereichweite eines Geräts der Gerätedichte in dessen Umgebung angepasst wird. Ein günstiger Wert für k ist daher unabhängig von der Gerätedichte. Eine günstige einheitliche Sendeleistung, die von allen Geräten verwendet werden soll, hängt hingegen von der Gerätedichte ab und müsste netzweit bestimmt und propagiert werden.

Ein weiterer Vorteil der Nearest-Neighbours-Topologien gegenüber anderen dezentral konstruierbaren Topologien wurde bereits in Abschnitt 4.2 beschrieben. Aus praktischen Gründen ist das sinnvollste Distanzmaß die Dämpfung, die ein Funksignal entlang eines Links erfährt, denn nur dann, wenn diese Dämpfung vergleichsweise gering ist, kann eine Übertragung mit geringer Sendeleistung erfolgreich sein. Ein Gerät kann durch das Erhöhen der Sendeleistung nur Links hinzugewinnen, die mit einer schlechteren Übertragungsqualität verbunden sind, die dem genannten Distanzmaß zufolge also mit einer größeren Distanz verbunden sind. Verfügt ein Gerät bereits über eine Mindestzahl an Links, so folgt daraus, dass Geräte, zu denen es durch eine Erhöhung der Sendeleistung neue Links etablieren könnte, nicht zu seinen nächsten Nachbarn gehören. Die Kenntnis dieser möglichen Links ist daher nicht erforderlich. Um andere Ansätze zu implementieren, z.B. die Angleichung der Topologie an einen Relative-Neighbourhood- oder Gabriel-Graph, müssen den Geräten hingegen alle möglichen Links und deren Qualität bekannt sein. Dazu müssen jedoch in regelmäßigen Abständen Signalisierungsnachrichten mit maximaler Leistung versendet werden.

Der Nachteil des Nearest-Neighbours-Ansatzes ist die in Abschnitt 4.2.1.2 beschriebene Schwäche, dass die resultierenden Topologien unter Umständen eine zu geringe Konnekti-

- + Hohe Kapazität.
- + Dezentrale Implementierung ausschließlich mit lokaler Information.
- + Keine Erkundung der lokalen Umgebung mit unnötig hohen Sendeleistungen.
- Maximale Konnektivität kann nicht garantiert werden.

Tabelle 7.1: Vor- und Nachteile des Nearest-Neighbours-Ansatzes

vität aufweisen. Die Wahrscheinlichkeit dafür sinkt mit steigender Mindestnachbarzahl jedoch rapide. Erwähnenswert ist in diesem Zusammenhang auch die in [HM04b] beschriebene Feststellung, dass das Lognormal-Shadowing (siehe Abschnitt 2.1.3.1) der Konnektivität drahtloser Netze zuträglich. Letztendlich muss aber anerkannt werden, dass hier ein Konflikt zwischen zwei Zielen vorliegt, nämlich der Kapazitätssteigerung und der Wahrung der Konnektivität. Werden die Sendeleistungen reduziert, um die Kapazität zu steigern, kann dies mit einer gewissen Wahrscheinlichkeit zu verminderter Konnektivität führen. Deshalb sollte in solchen Szenarien auf Sendeleistungsregelung verzichtet werden, in denen die Konnektivität wichtiger ist als die Kapazität, beispielsweise in Katastrophenszenarien, in denen eine geringe Menge außerordentlich wichtiger Daten ausgetauscht werden soll. Für ein offenes Netz (wie in Abschnitt 3.1 beschrieben) kann dagegen argumentiert werden, dass es sinnvoller ist, Einschränkungen der Konnektivität zu riskieren, als eine geringe Netzkapazität in Kauf zu nehmen.

Die Vor- und Nachteile der Nearest-Neighbours-Strategie sind in Tabelle 7.1 zusammengefasst. Es handelt sich um einen vielversprechenden Ansatz zur verteilten Sendeleistungsregelung, der aber bisher nicht verfolgt worden ist (außer in den im Rahmen der vorliegenden Arbeit entstandenen Veröffentlichungen [GdWMJ03] und [GdWMJ05]). Lediglich das Max-Degree-Verfahren, das in dem Sinne ähnlich ist, dass es ebenfalls die Nachbarzahl jedes einzelnen Geräts innerhalb eines bestimmten Bereichs zu halten versucht, ist bereits dezentral umgesetzt worden. Wie im vorangegangenen Kapitel gezeigt wurde, haben Max-Degree-Topologien jedoch deutliche Nachteile gegenüber Nearest-Neighbours-Topologien.

7.2 Arbeitsweise des Algorithmus

Das im Folgenden vorgestellte Verfahren zur dezentralen Sendeleistungsregelung versucht wie existierende Max-Degree-Verfahren die Nachbarzahl jedes Geräts innerhalb eines Zielbereichs $[k_{\min}; k_{\max}]$ zu halten, sorgt aber dafür, dass die Geräte im Sinne des Nearest-Neighbours-Ansatzes miteinander kooperieren und wurde deshalb mit der Bezeichnung *Cooperative Nearest Neighbours Topology Control* versehen. Es handelt sich um eine Weiterentwicklung des in [GdWMJ03], [GdWMJ05] präsentierten Verfahrens, dessen Grundidee auch schon in anderen Forschungsarbeiten aufgegriffen wurde ([BCF04]).

Es wird angenommen, dass die Hardware eine diskrete Menge möglicher Sendeleistungen vorgibt, so dass die Sendeleistung nicht auf jeden beliebigen Wert im Intervall zwischen einer Mindest- und einer Höchstleistung eingestellt werden kann, wie in einigen verwandten Arbei-

ten angenommen wird. Diese Annahme ist durch den Verweis auf aktuelle Technik gerechtfertigt. Beispielsweise unterstützen GSM-Endgeräte je nach Frequenzband und Geräteklasse bis zu 19 verschiedene Leistungsstufen [Eur99], und derzeit verfügbare WLAN-Hardware gar nur sechs Stufen (z.B. die Produktserie „Cisco Aironet 350“ [Cis]).

Das primäre Ziel des im Folgenden vorgestellten Algorithmus ist es, jedem Gerät mindestens $k_{\min}$ bidirektionale Links zur Verfügung zu stellen. Dies ist die Grundidee des Nearest-Neighbours-Ansatzes und durch eine geeignete Wahl des Werts $k_{\min}$ kann die Wahrscheinlichkeit dafür, dass das Netz unnötigerweise partitioniert wird, gering gehalten werden. Ein sekundäres Ziel ist es, dass die Anzahl bidirektionaler Links einen weiteren Schwellwert $k_{\max}$ nicht übersteigt. Da $k_{\min} = k_{\max}$ gewählt werden kann, ist dies lediglich eine Verallgemeinerung des Nearest-Neighbours-Ansatzes. Aus praktischen Gründen ist es jedoch sinnvoll, $k_{\max} > k_{\min}$ zu wählen, um eine unnötig hohe Anzahl von Sendeleistungsanpassungen zu vermeiden und in statischen Szenarien die Konvergenz hin zu einer stabilen Zuweisung von Sendeleistungen zu ermöglichen. Eine sinnvolle Spannbreite des Zielbereichs der Nachbarzahl $k_{\max} - k_{\min}$ hängt von der Granularität der möglichen Sendeleistungen ab; je größer die Einstellmöglichkeiten sind, desto größer sollte diese Spanne sein.

Der Algorithmus ist bereits von sich aus so beschaffen, dass eine unterschiedliche Wahl der Parameter $k_{\min}$ und $k_{\max}$ keine negativen Auswirkungen hat. Es ist sogar durchaus sinnvoll, die Differenz dieser Parameter individuell an die Granularität der Einstellmöglichkeiten für die Sendeleistung eines Geräts anzupassen. Beim Entwurf der im Folgenden beschriebenen Mechanismen ist außerdem darauf geachtet worden, dass sie auch dann effektiv sind, wenn die Geräte verschiedene Parameterbelegungen verwenden. Wie in Abschnitt 3.3 dargelegt wurde, ist dies wichtig, um auch praktisch einen dezentralen Betrieb zu ermöglichen.

Die Signalisierung geschieht ausschließlich durch Beacons, die jedes Gerät in regelmäßigen Abständen als Broadcast-Rahmen versendet. Ein solcher Mechanismus kann auch ohne Sendeleistungsregelung verwendet werden, um zuverlässige von unzuverlässigen Links zu unterscheiden (siehe Abschnitt 2.2.2). Wenn Routingprotokolle für drahtlose Multihop-Netze in vielen Fällen auch darauf ausgelegt sind, dass die Geräte keinerlei Kenntnis ihrer lokalen Umgebungen haben, so zeigen später präsentierte Simulationsergebnisse, dass allein aus der Einschränkung des Routings auf Links, die als zuverlässig identifiziert wurden, eine deutliche Leistungssteigerung resultiert.

Ein Mechanismus zur Linkerkennung ist ebenfalls notwendig, wenn Multirate-Hardware im Sinne der Kapazität des Netzes möglichst gut genutzt werden soll, was bisher eine kaum beachtete Herausforderung drahtloser Multihop-Netze ist. Werden Pfade nämlich wie in Abschnitt 2.3.4 beschrieben durch das Fluten des Netzes mit Broadcast-Rahmen etabliert, so hat deren Datenrate einen wesentlichen Einfluss darauf, wie Kommunikationspfade beschaffen sind. Broadcast-Rahmen, die mit einer geringeren Datenrate versendet werden, also robust kodiert sind, können auch über Links empfangen werden, auf denen das Signal eine stärkere Dämpfung erfährt und die somit nicht für hohe Datenraten geeignet sind. Dies bewirkt aber, dass sowohl die Kapazität entlang des Pfads als auch die Gesamtkapazität des Netzes unter Umständen unnötigerweise verringert wird (Letzteres aus dem Grund, dass Übertragungen

mit geringeren Datenraten den gemeinsam genutzten Übertragungskanal länger blockieren). Daraus sollte aber nicht die Konsequenz gezogen werden, dass die zur Routensuche eingesetzten Broadcast-Rahmen mit hohen Datenraten verschickt werden sollten, denn auf diese Weise können Routen, die in weniger dichten Netzen zwangsläufig auch Links mit stärkerer Dämpfung enthalten müssen, gar nicht zustande kommen.

Ein Beacon-Mechanismus, wie er hier verwendet wird, um zuverlässige, bidirektionale Links zu identifizieren, ist vermutlich auch eine geeignete Grundlage, um dieser Herausforderung beizukommen: Indem die Datenrate, mit der Beacons versendet werden, nicht festgelegt ist, sondern zwischen verschiedenen verfügbaren Werten wechselt, kann die Zuverlässigkeit von Links in Abhängigkeit der Datenrate ermittelt werden. Die maximale Datenrate, bei der ein Link als zuverlässig erachtet wird, kann dann in die Routingmetrik einfließen. Dies ist auch durch die einfache Erweiterung von reaktiven Routingprotokollen wie AODV möglich, wie in [GdWM04] ausgearbeitet wird. Die genaue Definition der Routingmetrik hätte für diese Lösung einen wesentlichen Einfluss auf die Netzkapazität und bedarf somit einer eingehenden Untersuchung, die über den Rahmen der vorliegenden Arbeit hinausgeht.

Ein Beacon-Mechanismus zur Linkerkennung kann schließlich auch zur statistischen Auswertung von Linklebensdauern und damit zur Identifikation von Links genutzt werden, die trotz der Mobilität der Endgeräte mit einer gewissen Wahrscheinlichkeit noch einige Zeit verfügbar sein werden [GdWFM02]. Aus all diesen Gründen muss ein Mechanismus wie der hier beschriebene in drahtlosen Multihop-Netzen als unverzichtbar angesehen werden. Die zur Sendeleistungsregelung benötigte Signalisierung wird den Beacons hinzugefügt und vergrößert sie lediglich um eine geringe, konstante Anzahl von Bits.

7.3 Sammeln lokaler Topologieinformation

Mit Hilfe eines Beacon-Mechanismus können bestehende Links erkannt und über ihre Zuverlässigkeit entschieden werden. Wie in Abschnitt 2.2.2 beschrieben, kann ein Gerät durch den Empfang von Beacons zunächst feststellen, von welchen anderen Geräten zuverlässige unidirektionale Links zu ihm selbst ausgehen. Indem deren Adressen als *Echos* in die gesendeten Beacons eingefügt werden, können zuverlässige bidirektionale Links von einem Gerät daran erkannt werden, dass Beacons über einen (zuverlässigen) unidirektionalen Link empfangen werden, die seine eigene Adresse enthalten.

Nach welchen konkreten Kriterien die Zuverlässigkeit eines Links entschieden wird, ist für den hier vorgestellten Algorithmus nicht relevant. Die zur Evaluierung verwendete Implementierung bewertet einen unidirektionalen Link als zuverlässig, wenn der Anteil der erfolgreich empfangenen unter den letzten l_{win} versendeten Beacons mindestens l_{up} erreicht und blendet ihn wieder aus, wenn diese Zahl unter l_{down} fällt. Ein hinreichend großer Abstand $l_{\text{up}} - l_{\text{down}}$ ist sinnvoll, um häufige Änderungen von Linkzuständen zu vermeiden (ein vergleichbarer Mechanismus, die Verwendung eines *Hysterese-Werts*, verhindert in GSM-Systemen, dass ein mobiles Endgerät unnötig häufig die Location-Area wechselt, wenn zwischen zwei Basissta-

tionen eine vergleichbar gute Verbindung besteht [Eur98]).

Abbildung 7.2 stellt den beschriebenen Mechanismus beim Beacon-Empfang in Pseudocode-Schreibweise dar. Wie dort zu sehen ist, enthalten die Beacons auch das Beacon-Intervall des Senders, damit die Geräte in der Lage sind, brauchbare Timeout-Werte für die Zuverlässigkeit von Links zu berechnen. Außerdem zeigt die Abbildung, welche Informationen noch ausgetauscht werden und wie verschiedene Nachbarmengen verwaltet werden, um die Kooperation zwischen den Geräten zu ermöglichen. Dies wird allerdings erst in den nachfolgenden Abschnitten erläutert.

Um den durch das Beaconsing verursachten Overhead zu beschränken, sollte die maximale Anzahl der Echos pro Beacon begrenzt werden. Die zur Auswertung benutzte Implementierung schreibt maximal 64 Echos in ein Beacon. Dabei werden in erster Linie die Adressen der Nachbarn weggelassen, die mehr als $k_{\max}$ Nachbarn haben oder die signalisiert haben, dass der Link redundant ist (diese Signalisierung wird in Abschnitt 7.5 eingeführt). Um den Pseudocode übersichtlich zu halten, ist diese Begrenzung in Abbildung 7.3 nicht enthalten.

Werden IPv4-Adressen verwendet, resultiert der beschriebene Mechanismus in einem Overhead von 32 Bits pro Echo pro Beacon. Bei einer Obergrenze von 64 Echos pro Beacon beträgt das maximale Datenaufkommen für Echos damit 256 Bytes pro Beacon, was ein vertretbarer Aufwand ist. Handelt es sich hingegen um ein IPv6-Netz, ist fraglich, ob der Protokolloverhead sich in sinnvollen Grenzen hält. Grundsätzlich lässt sich aber die gleiche Funktionalität erreichen, indem Beacons nicht durch die Absenderadresse, sondern durch Zufallswerte identifiziert werden. Wählen zwei Geräte in einem bestimmten Zeitraum den gleichen Wert, so kann dies zwar in einer fehlerhaften Sicht der Netztopologie resultieren, da der Anteil der erfolgreich empfangenen Beacons dieser Geräte von anderen Netzteilnehmern nicht mehr richtig berechnet werden kann. Stehen für die Zufallswerte aber genügend Bits zur Verfügung, dann treten solche Ereignisse so selten ein, dass sie langfristig gesehen keinen nennenswerten Einfluss auf die Netztopologie haben.

Es ist noch anzumerken, dass bei der Verwendung fester Werte für l_{up} und l_{down} im Zusammenspiel mit ungesicherten Beacon-Broadcasts die Rate der erfolgreich über einen Link empfangenen Beacons von der allgemeinen Last in der Umgebung des Links abhängt. Deshalb werden zunächst als zuverlässig erachtete Links durch das Vorhandensein zusätzlicher Last unter Umständen ausgeblendet. Dies kann auf der einen Seite durchaus sinnvoll sein, da die Fehlerrate ungesicherter Broadcasts auch Rückschlüsse auf die Anzahl der Paketwiederholungen gesicherter Datenübertragungen zulässt. Bei hoher Last hat dies allerdings zur Folge, dass einige Geräte ihre Sendeleistungen erhöhen, um vermeintlich verloren gegangene Links wieder aufzubauen oder durch neue zu ersetzen. Aus diesem Grund könnte es sinnvoll sein, die Werte l_{up} und l_{down} dynamisch an die Last anzupassen, die in der Umgebung beobachtet wird. Eine tiefergehende Untersuchung von Kriterien zur Evaluierung der Zuverlässigkeit von Links ist jedoch außerhalb des Rahmens dieser Arbeit.

7.4 Erhöhen der Sendeleistung

Es gibt zwei Maßnahmen, die einem Gerät A zu mehr bidirektionalen Links verhelfen können. Zum einen können Geräte, die sich bereits in der Sendereichweite von A befinden, ihre Sendeleistung erhöhen, falls sich A noch nicht in ihrer Sendereichweite befindet. Dies erkennen sie daran, dass sie eine hinreichende Anzahl der letzten Beacons von A empfangen haben, die ihre eigene Adresse aber nicht enthalten. Dadurch werden unidirektionale Links zu bidirektionalen Links ausgebaut. Zum anderen kann das Gerät A selbst seine Sendeleistung erhöhen, um weitere Links zu Geräten in größerer Distanz zu etablieren. An A eingehende unidirektionale Links können durch das Erhöhen der Sendeleistung von A zu bidirektionalen Links ausgebaut werden, oder es entstehen ausgehende unidirektionale Links, die wiederum zu bidirektionalen Links ausgebaut werden können, indem andere Geräte ihre Sendeleistungen erhöhen.

Hat ein Gerät eine zu geringe Anzahl bidirektionaler Links, so signalisiert es dies in seinen Beacons mit Hilfe des dafür vorgesehenen `help`-Felds. Der Empfänger eines Beacons mit `help` > 0 erhöht seine Sendeleistung unabhängig davon, wie viele bidirektionale Links er bereits hat, wenn das Beacon über einen unidirektionalen Link empfangen wurde und wenn er dies nicht bereits auf ein vorangegangenes Beacon desselben Geräts mit demselben Wert im `help`-Feld hin getan hat. Ein Gerät mit $k < k_{\min}$ Nachbarn versendet zunächst mehrere Beacons mit dem gleichen `help`-Wert, so dass daraufhin neu etablierte bidirektionale Links auch erkannt werden können. Falls dann immer noch $k < k_{\min}$ ist, kann der Vorgang mit neuen `help`-Werten wiederholt werden, da potenzielle neue Nachbarn unter Umständen mehrmals ihre Sendeleistung inkrementieren müssen, damit ein bidirektionaler Link aufgebaut werden kann.

Verfügt das Gerät dann immer noch über zu wenig bidirektionale Links, so erhöht es die eigene Sendeleistung. Dadurch, dass zunächst unidirektionale Links zu bidirektionalen Links ausgebaut werden, ist dafür gesorgt, dass Verbindungen möglichst zwischen nächsten Nachbarn aufgebaut werden. Wie bereits in Abschnitt 4.2.1.2 erläutert, kann ein Gerät durch das Erhöhen der eigenen Sendeleistung nur Links zu Nachbarn aufbauen, die weiter entfernt sind als jene, die sich zuvor bereits in seiner Sendereichweite befanden (wobei die Entfernung nicht durch die räumliche Distanz, sondern durch die Dämpfung von Funksignalen entlang des Links gegeben ist). Auf diese Weise wird beispielsweise die in Abbildung 4.2 dargestellte Situation sinnvoll behandelt, wozu die Max-Degree-Verfahren nicht in der Lage sind.

Nach einer Erhöhung der eigenen Sendeleistung wird `help` zunächst wieder auf 0 gesetzt, da zunächst überprüft werden soll, ob weitere bidirektionale Links etabliert werden könnten, ohne dass andere Geräte ihre Sendeleistungen erhöhen müssen.

Das beschriebene Verhalten ist in Abbildung 7.4 in Form von Pseudocode dargestellt, der unmittelbar vor Versenden eines Beacons ausgeführt wird, wenn $k < k_{\min}$ ist. Alle verwendeten Variablen sind in globalem Kontext gültig und werden teilweise auch im Pseudocode für das Gerüst des Algorithmus in Abbildung 7.7 verwendet.

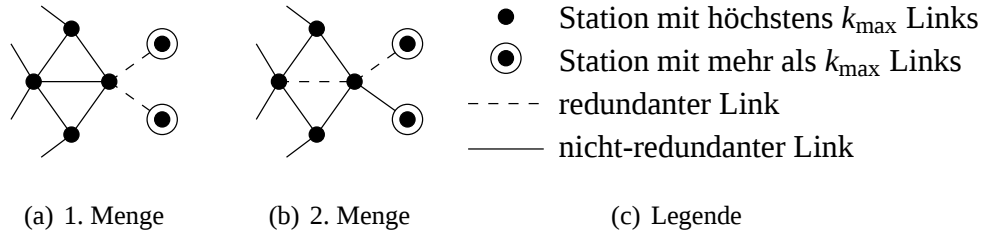


Abbildung 7.1: Zwei verschiedene Mengen redundanter Links eines Geräts ($k_{\min} = 3$)

7.5 Verringern der Sendeleistung

Ein Gerät, das mehr als $k_{\max}$ bidirektionale Links hat, darf nicht ohne weiteres seine Sendeleistung verringern, da seine Nachbarn auf die betroffenen Links angewiesen sein könnten. Dies ist beispielsweise dann der Fall, wenn ein Nachbar gerade $k_{\min}$ Links hat, denn dann muss jeder einzelne von ihnen aufrechterhalten werden, damit diese Mindestnachbarzahl nicht unterschritten wird. Hat ein Gerät $k > k_{\min}$ Links und $l > k - k_{\min}$ Nachbarn mit mehr als $k_{\max}$ Links, dann können zwar einige, aber nicht alle dieser Nachbarn ihre Sendeleistungen verringern. Die Schwierigkeit liegt darin, zu verhindern, dass mehr als $k - k_{\min}$ Nachbarn dieses Geräts ihre Sendeleistungen zeitnah verringern. Andernfalls blieben diesem Gerät weniger als $k_{\min}$ Links, und auf die daraufhin versendeten `help`-Beacons müssten die kurz zuvor verringerten Sendeleistungen wieder erhöht werden. Dies würde sich beliebig oft wiederholen, sofern alle Geräte mit mehr als $k_{\max}$ Links zeitnah nach dem ersten Empfang eines Beacons mit gelöschtem `help`-Flag ihre Sendeleistungen reduzieren. Eine mögliche Lösung ist ein exponentieller Backoff, der dafür sorgt, dass immer längere Wartezeiten zwischen den Sendeleistungsverringerungen eingehalten werden, die wieder rückgängig gemacht werden müssen. Diese Idee wird in Abschnitt 7.5.3 genauer ausgeführt.

Im Allgemeinen ist also nicht eindeutig, auf welche Links ein Gerät angewiesen ist und auf welche nicht. Insbesondere können einzelne Links hierfür nicht isoliert betrachtet werden, sondern dies muss im Kontext mehrerer Links geschehen. Deshalb wird eine Menge $R \subset L$ von Links als *redundant* bezeichnet, wenn jedes Gerät v , das mit einem oder mehreren dieser Links inzident ist, mehr als $k_{\max}$ Links hat oder mindestens $k_{\min}$ Links hat, die nicht in R sind:

$$\forall v \in S : (\exists w \in S : (v, w) \in R) \Rightarrow (\deg(v) > k_{\max} \vee |\{w \in S : (v, w) \in L \setminus R\}| \geq k_{\min})$$

Geräte, die mindestens $k_{\min}$ und höchstens $k_{\max}$ Links haben, behalten also mindestens $k_{\min}$ Links, wenn eine redundante Menge von Links aus der Netztopologie entfernt wird. Für Geräte mit mehr als $k_{\max}$ Links kann dies hingegen nicht garantiert werden, wenn die Sendeleistungen dieser Geräte verringert werden. Daher ist es wichtig, dass die Differenz $k_{\max} - k_{\min}$ so gewählt wird, dass nur mit geringer Wahrscheinlichkeit mehr als $k_{\max} - k_{\min}$ Links durch die Verringerung der Sendeleistung um eine Stufe verloren gehen.

Abbildung 7.1 zeigt beispielhaft zwei mögliche Mengen redundanter Links eines Geräts. Da

das Gerät 5 Links hat und $k_{\min} = 3$ ist, kann eine Menge redundanter Links dieses Geräts höchstens zwei Elemente umfassen, die dargestellten Mengen sind also maximal.

Der verteilte Algorithmus sieht zunächst vor, dass Geräte signalisieren müssen, wenn sie mehr als $k_{\max}$ Links haben. Im Folgenden werden verschiedene Mechanismen zur Verringerung der Sendeleistung eingeführt, die in Abbildung 7.5 in Pseudocode-Notation zusammengefasst sind.

7.5.1 Die einfache Variante

Eine einfache Variante des Algorithmus wurde in [GdWMJ03], [GdWMJ05] vorgeschlagen. Hier signalisieren Geräte mit mehr als $k_{\max}$ Links ihren Willen zur Verringerung der Sendeleistung durch das Setzen eines `satisfied`-Flags, und sie verringern ihre Sendeleistung nur dann, wenn dieses Flag in den Beacons aller Nachbarn ebenfalls gesetzt ist. Dies ist offensichtlich eine sehr restriktive Bedingung, denn ein Gerät darf seine Sendeleistung beispielsweise schon nicht verringern, wenn ein einziger Nachbar weniger als $k_{\max}$, aber mehr als $k_{\min}$ Links hat.

7.5.2 Signalisierung redundanter Links

In der hier beschriebenen Variante des Algorithmus wird deshalb eine zusätzliche Signalisierung verwendet, um redundante Links anzuzeigen. Dies kann effizient realisiert werden, indem die Adress-Echos in zwei getrennte Listen eingeteilt werden (im Pseudocode als $E_{\text{nonredundant}}$ und $E_{\text{redundant}}$ bezeichnet). Ein Gerät mit $k > k_{\min}$ Links kann so bis zu $k - k_{\min}$ Nachbarn mit mehr als $k_{\max}$ Links anzeigen, dass dessen Links zu ihnen redundant sind. Falls sie mehr als $k - k_{\min}$ solcher Nachbarn hat, müssen $k - k_{\min}$ von ihnen ausgewählt werden. Da nur Links zu Geräten mit mehr als $k_{\max}$ Nachbarn als redundant gekennzeichnet werden, tritt eine Situation wie in Abbildung 7.1(b) nicht auf: Den Link zwischen den beiden Geräten mit Knotengrad 5 in die Menge redundanter Links aufzunehmen ist deshalb nicht sinnvoll, weil es gar nicht beabsichtigt ist, die Sendeleistung eines dieser Geräte zu senken.

Die Auswahl redundanter Links wird mit der Zeit verändert, um allen Geräten, die ihre Absicht zur Verringerung der Sendeleistung anzeigen, langfristig auch die Möglichkeit dazu zu geben. Andernfalls könnten Links beliebig lange ohne Wirkung als redundant klassifiziert sein, weil die Möglichkeit zur Verringerung der Sendeleistung von den entsprechenden Geräten nicht genutzt werden kann. Da andererseits das Wegfallen von Links als Folge von verringerten Sendeleistungen zunächst festgestellt werden muss, ist es auch nicht sinnvoll, diese Auswahl zu rasch zu verändern. Deshalb ist die Dynamik dieser Auswahl dadurch implementiert, dass ein redundanter Link in jedem Beacon-Intervall mit einer bestimmten, relativ geringen Wahrscheinlichkeit als nicht redundant klassifiziert wird, woraufhin die Auswahl redundanter Links teilweise erneuert werden kann. Diese Wahrscheinlichkeit ist im Pseudocode mit $p_{\text{redundancy_fluctuation}}$ bezeichnet.

Da ein Gerät mit mehr als $k_{\max}$ Nachbarn warten muss, bis diese Nachbarn alle gleichzeitig gerade ihm eine Verringerung der Sendeleistung erlauben, erscheint es in dichteren Netzen sinnvoll, die Auswahl redundanter Links zu fokussieren. Dazu können Geräte mit höheren $k - k_{\max}$ bevorzugt werden, wodurch gerade dort eine Verringerung der Sendeleistung angestrebt wird, wo die Konkurrenz um das geteilte Übertragungsmedium besonders hoch ist. Deshalb reicht ein einfaches Flag zur Signalisierung nicht aus; stattdessen setzen Geräte mit $k > k_{\max}$ in ihren Beacons `satisfied = k - k_{\max}`. Diese Variante ist die Standardvariante für die Auswertungen in Kapitel 8.

7.5.3 Versuchsweises Reduzieren

Der im vorangegangenen Abschnitt beschriebene Mechanismus ist sehr vorsichtig: Sendeleistungen werden grundsätzlich nur dann verringert, wenn sichergestellt ist, dass dadurch kein Nachbar einen nicht redundanten Link verlieren könnte, so dass ein einziges Gerät mit $k = k_{\min}$ alle seine Nachbarn blockiert. Dabei kann ein Gerät, das auch an nicht redundanten Links beteiligt ist, seine Sendeleistung durchaus verringern, wenn dadurch nur redundante Links wegfallen.

Ob ein Link von einer Sendeleistungsverringern betroffen sein könnte oder nicht, ließe sich unter Umständen durch Verfahren abschätzen, die auf Signalstärkemessungen beruhen. Wegen der komplizierten Eigenschaften der Ausbreitung von Radiosignalen und der daraus resultierenden starken Schwankungen, denen die Empfangsstärke eines Signals unterworfen ist (siehe Abschnitt 2.1.3), ist es praktisch jedoch sehr schwierig, ein solches Verfahren in geeigneter Weise zu realisieren. Deshalb geht das hier vorgestellte Verfahren einen einfacheren Weg: Geräte mit $k > k_{\max}$ Nachbarn verringern *versuchsweise* ihre Sendeleistung. Gehen dadurch nicht redundante Links verloren, stellen sie dies dadurch fest, dass ehemalige Nachbarn über einen unidirektionalen Link `help` signalisieren, und die Sendeleistung muss wieder erhöht werden. In diesem Fall sorgt ein exponentieller Backoff dafür, dass die Anzahl der unnötigen Sendeleistungsanpassungen sich in akzeptablen Grenzen hält. Dazu wird für jeden Nachbarn n ein Timeout-Wert `freezen` gespeichert. Eine versuchsweise Verringerung der Sendeleistung erfolgt nur dann, wenn dieser Timeout für alle Nachbarn abgelaufen ist. Falls der ehemalige Nachbar innerhalb eines gewissen Zeitraums nach einer Verringerung der Sendeleistung über einen inzwischen unidirektionalen Link `help` signalisiert, wird die Zeitspanne des Timeouts für diesen Nachbarn `timeoutn` verdoppelt. Die Länge dieses Zeitraums ist mit dem doppelten l_{win} des Nachbarn bewusst etwas länger gewählt, weil eine Sendeleistungsverringern nicht sofort zum Verschwinden des Links führen muss, sondern dessen Qualität auch so verschlechtern kann, dass sie grenzwertig in Bezug auf die vorgegebenen Linkerkennungsparameter ist, so dass der Link erst etwas später abbricht.

Die Verwaltung eines Timeout-Werts für jeden Nachbarn mag als vergleichsweise großer Aufwand erscheinen, der durch die Verwendung eines einzelnen globalen Timeout-Werts deutlich reduziert werden könnte. Es ist aber zu beachten, dass es sich lediglich um einen Zahlenwert handelt, der pro Nachbar gespeichert werden muss, dass aber keine Timer-Mechanismen ge-

cooperation, redundancy_signalling, trial_reduction	Auswahl einer konkreten Variante des Algorithmus Wertebereich: {enabled, disabled}
$[k_{\min}; k_{\max}]$	Zielbereich der Nachbarzahl
t_{beacon}	Maximaler Zeitabstand zwischen aufeinander folgenden Beacons
t_{jitter}	Maximale Verkürzung des Beacon-Intervalls
myaddr	Adresse des lokalen Geräts
$t_{\text{powerdown}}$	Mindestzeit zwischen zwei aufeinander folgenden Sendeleistungsverringerungen
n_{waitlink}	Anzahl der Wiederholungen eines help-Werts
n_{powerup}	Anzahl der unterschiedlichen help-Werte bis zur Erhöhung der eigenen Sendeleistung
$t_{\text{maxtimeout}}$	Maximale Zeitspanne zwischen sukzessiven versuchsweisen Sendeleistungsverringerungen
$l_{\text{win}}, l_{\text{up}}, l_{\text{down}}$	Linkerkennungsparameter
$p_{\text{redundancy_fluctuation}}$	Wahrscheinlichkeit, mit der ein Link seine Redundanz-Markierung verliert

Tabelle 7.2: Vom Algorithmus verwendete Konstanten

nutzt werden. Ein Gerät, das seine Sendeleistung reduzieren möchte, muss lediglich überprüfen, ob das Maximum all dieser Timeout-Werte vergangen ist oder nicht.

Bei der Verwaltung getrennter Timeout-Werte für verschiedene Links ist jeder dieser Werte genau auf die Zeitspanne angepasst, während der die Beibehaltung der aktuellen Sendeleistung zur Aufrechterhaltung des zugehörigen Links erforderlich war. Geht ein Link aus anderen Gründen verloren, beispielsweise weil ein Gerät ausgeschaltet oder fortbewegt wurde, können die anderen Timeout-Werte einfach weiterverwendet werden. Bei Verwaltung eines globalen Timeout-Werts hingegen müsste zunächst festgelegt werden, wie dieser Wert in einer solchen Situation zu behandeln ist, in jedem Fall können sich jedoch Nachteile ergeben: Wird der Timeout-Wert in diesem Fall grundsätzlich nicht zurückgesetzt, so wird vielleicht unnötig lange eine Sendeleistung beibehalten, die für einen Link benötigt wurde, der gar nicht mehr vorhanden ist. Wird der Timeout-Wert in diesem Fall dagegen zurückgesetzt, muss er sich mit Hilfe des exponentiellen Backoffs erneut auf einen brauchbaren Wert einpendeln, was zu einer unnötig hohen Zahl von Sendeleistungsanpassungen führt, die gemäß den Überlegungen in Abschnitt 3.3 vermieden werden sollte.

7.6 Der Algorithmus im Überblick

Abbildung 7.7 zeigt das Grundgerüst für verschiedene Varianten des Algorithmus, und die im Pseudocode verwendeten Konstanten sind in Tabelle 7.2 zusammengefasst. Die Initialisierung

globaler Variablen wird in Abbildung 7.6 durchgeführt.

Im Laufe dieses Kapitels sind verschiedene Mechanismen erläutert worden, die in ihrer Gesamtheit die Standardvariante des Verfahrens ausmachen. Im folgenden Kapitel wird der Frage nachgegangen, ob und unter welchen Umständen diese im einzelnen sinnvoll sind. Dazu werden unterschiedliche Varianten des Verfahrens untersucht, die im abgebildeten Pseudocode durch die Belegung der drei Variablen „cooperation“, „redundancy_signalling“ und „trial_reduction“ ausgewählt werden, die in der Standardvariante allesamt auf „enabled“ gesetzt sind.

Zunächst stellt sich die Frage, welche Vorteile der Nearest-Neighbours-Ansatz gegenüber dem Max-Degree-Ansatz bietet, der von der Variante mit deaktivierter Kooperation (`cooperation = disabled`) implementiert wird.

Die Standardvariante enthält weiterhin mehrere Mechanismen zur Sendeleistungsreduktion: Die Signalisierung redundanter Links und das versuchsweise Reduzieren mit exponentiellem Backoff. Deshalb wird untersucht, wie gut jeder dieser Mechanismen alleine unter verschiedenen Umständen funktioniert, es werden also eine Variante mit deaktivierter versuchsweiser Reduktion (`trial_reduction = disabled`) und eine mit deaktivierter Redundanzsignalisierung (`redundancy_signalling = disabled`) betrachtet. Da beide Mechanismen wiederum deutlich komplizierter sind als der einfache Reduktionsmechanismus in [GdWMJ03] und [GdWMJ05], wird auch diese einfache Variante in die folgenden Untersuchungen einbezogen (`trial_reduction = redundancy_signalling = disabled`).

Abbildung 7.2 : Prozedur `recv_beacon`

Beacon-Felder : n : Absender-Adresse
 s : Sequenznummer
 beacon_interval : t_{beacon} des Senders
 detection_window : l_{win} des Senders
 $E_{\text{redundant}}$: Adress-Echos für als redundant eingestufte Nachbarn
 $E_{\text{nonredundant}}$: Adress-Echos für alle anderen Nachbarn
 help : Inhalt des help -Sequenznummer
 satisfied : Anzahl überzähliger Links

```

 $B_n \leftarrow \{s' \in B_n : s' > s - l_{\text{win}}\} \cup \{s\}$ 
if  $|B_n|/l_{\text{win}} \geq l_{\text{up}} \wedge n \notin N_{\text{unidir}}$  then
     $N_{\text{unidir}} \leftarrow N_{\text{unidir}} \cup \{n\}$ 
    // Globale Variablen (re-)initialisieren
     $\text{lasthelp}_n \leftarrow 0$ 
     $\text{unsure}_n \leftarrow 0$ 
     $\text{timeout}_n \leftarrow 0$ 
     $\text{redundant}_n \leftarrow \text{false}$ 
     $w_n \leftarrow \text{detection\_window}$ 
if  $n \in N_{\text{unidir}}$  then
     $\text{expiry\_seqno} \leftarrow \min \{s' \in B_n : |\{s'' \in B_n : s'' > s'\}|/l_{\text{win}} < l_{\text{down}}\}$ 
     $\text{expiry\_time} \leftarrow t_{\text{now}} + (\text{expiry\_seqno} - s + l_{\text{win}} + 0,5) \cdot \text{beacon\_interval}$ 
     $\text{restart\_timer}(\text{expiration\_timer}_n, \text{expiry\_time})$ 
    // Bei Ablauf dieses Timers wird  $n$  aus allen Nachbarlisten entfernt.
    if  $\text{myaddr} \in E_{\text{redundant}} \cup E_{\text{nonredundant}}$  then
         $N_{\text{bidir}} \leftarrow N_{\text{bidir}} \cup \{n\}$ 
        if  $\text{myaddr} \in E_{\text{redundant}}$  then
             $N_{\text{redundant}} \leftarrow N_{\text{redundant}} \cup \{n\}$ 
        else
             $N_{\text{redundant}} \leftarrow N_{\text{redundant}} \setminus \{n\}$ 
    else
         $N_{\text{bidir}} \leftarrow N_{\text{bidir}} \setminus \{n\}$ 
         $N_{\text{redundant}} \leftarrow N_{\text{redundant}} \setminus \{n\}$ 
    if  $\text{help} > \text{lasthelp}_n$  then
         $N_{\text{help}} \leftarrow N_{\text{help}} \cup \{n\}$ 
     $\text{lasthelp}_n \leftarrow \text{help}$ 
    if  $\text{satisfied} > 0$  then
         $N_{\text{satisfied}} \leftarrow N_{\text{satisfied}} \cup \{n\}$ 
    else
         $N_{\text{satisfied}} \leftarrow N_{\text{satisfied}} \setminus \{n\}$ 

```

Abbildung 7.3 : Prozedur send_beacon

```

Parameter : help : help-Feld
              satisfied : satisfied-Feld

 $E_{\text{redundant}} \leftarrow \emptyset$ 
 $E_{\text{nonredundant}} \leftarrow \emptyset$ 
if redundancy_signalling = enabled then
  forall  $n \in N_{\text{unidir}} : \text{redundant}_n$  do
    if  $n \notin N_{\text{bidir}} \cap N_{\text{satisfied}}$  then
       $\text{redundant}_n \leftarrow \text{false}$ 
    else if  $\text{random\_uniform}([0; 1]) < p_{\text{redundancy\_fluctuation}}$  then
       $\text{redundant}_n \leftarrow \text{false}$ 

  while  $|\{n \in N_{\text{bidir}} : \neg \text{redundant}_n\}| > k_{\min} \wedge \exists n \in N_{\text{satisfied}} \cap N_{\text{bidir}} : \neg \text{redundant}_n$  do
     $n \leftarrow \text{zufälliges Element aus } \{n \in N_{\text{satisfied}} \cap N_{\text{bidir}} : \neg \text{redundant}_n\}$ 
     $\text{redundant}_n \leftarrow \text{true}$ 

  forall  $n \in N_{\text{unidir}}$  do
    if  $\text{redundant}_n$  then
       $E_{\text{redundant}} \leftarrow E_{\text{redundant}} \cup \{n\}$ 
    else
       $E_{\text{nonredundant}} \leftarrow E_{\text{nonredundant}} \cup \{n\}$ 

  else
     $E_{\text{nonredundant}} \leftarrow N_{\text{unidir}}$ 

  beacon  $\leftarrow (\text{myaddr}, \text{seqno}, t_{\text{beacon}}, l_{\text{win}}, E_{\text{redundant}}, E_{\text{nonredundant}}, \text{help}, \text{satisfied})$ 
  seqno  $\leftarrow \text{seqno} + 1$ 
  broadcast(beacon)

```

Abbildung 7.4 : Prozedur increase_k

```

help_repeat  $\leftarrow \text{help\_repeat} + 1$ 
if help_repeat =  $n_{\text{waitlink}}$  then
  help_repeat  $\leftarrow 0$ 
  help_count  $\leftarrow \text{help\_count} + 1$ 
  if cooperation = disabled  $\vee$  help_count =  $n_{\text{powerup}}$  then
    freeze  $\leftarrow \max\{\text{freeze}, t_{\text{now}} + t_{\text{powerdown}}\}$ 
    increase_power()
  else
    if help_count >  $n_{\text{powerup}}$  then
      help_count  $\leftarrow 0$ 
    help_seq  $\leftarrow \text{help\_seq} + 1$ 

if help_count <  $n_{\text{powerup}}$  then
  help  $\leftarrow \text{help\_seq}$ 
else
  help  $\leftarrow 0$ 
return help

```

Abbildung 7.5 : Prozedur decrease_k

```

if cooperation = disabled then
  | allow_reduction  $\leftarrow$  true
else if trial_reduction = enabled then
  | allow_reduction  $\leftarrow \forall n \in N_{\text{bidir}} \setminus N_{\text{redundant}} : \text{freeze}_n \geq t_{\text{now}}$ 
else if redundancy_signalling = enabled then
  | allow_reduction  $\leftarrow N_{\text{bidir}} \setminus N_{\text{redundant}} = \emptyset$ 
else
  | // Einfache Variante gemäß [GdWMJ03], [GdWMJ05]
  | allow_reduction  $\leftarrow N_{\text{bidir}} \setminus N_{\text{satisfied}} = \emptyset$ 
if allow_reduction then
  | decrease_power()
  | forall  $n \in N_{\text{bidir}}$  do
  |   |  $\text{unsure}_n \leftarrow t_{\text{now}} + 2 \cdot w_n$ 
  |   |  $\text{freeze} \leftarrow t_{\text{now}} + t_{\text{powerdown}}$ 
  |   |  $\text{satisfied} \leftarrow 0$ 
else
  |  $\text{satisfied} \leftarrow |N_{\text{bidir}}| - k_{\text{max}}$ 
return satisfied

```

Abbildung 7.6 : Prozedur initialise_global_variables

```

// Frühester Zeitpunkt der nächsten Sendeleistungsverringern
freeze  $\leftarrow -\infty$ 
// Zahl der versendeten helps
help_count  $\leftarrow 0$ 
// Zahl der Wiederholungen des aktuellen helps
help_repeat  $\leftarrow 0$ 
// Aktuelle help-Sequenznummer
help_seq  $\leftarrow 1$ 
// Aktuelle Beacon-Sequenznummer
seqno  $\leftarrow 0$ 
// Verschiedene Nachbarmengen
 $N_{\text{unidir}} \leftarrow \emptyset$ 
 $N_{\text{bidir}} \leftarrow \emptyset$ 
 $N_{\text{redundant}} \leftarrow \emptyset$ 
 $N_{\text{help}} \leftarrow \emptyset$ 
 $N_{\text{satisfied}} \leftarrow \emptyset$ 

```

Abbildung 7.7 : Grundgerüst des verteilten Algorithmus zur Sendeleistungsregelung

```

initialise_global_variables()
 $t_{\text{init}} \leftarrow t_{\text{now}} + 2 \cdot t_{\text{beacon}} \cdot l_{\text{win}} \cdot l_{\text{up}}$ 
while true do
   $\text{help} \leftarrow 0$ 
   $\text{satisfied} \leftarrow 0$ 
  if  $t_{\text{now}} > t_{\text{init}}$  then
     $k \leftarrow |N_{\text{bidir}}|$ 
    if  $\text{cooperation} = \text{enabled} \wedge N_{\text{help}} \setminus N_{\text{bidir}} \neq \emptyset$  then
      forall  $n \in N_{\text{help}} \setminus N_{\text{bidir}}$  do
        if  $\text{unsure}_n > t_{\text{now}}$  then
           $\text{timeout}_n \leftarrow \min\{t_{\text{maxtimeout}}, \min\{t_{\text{powerdown}}, 2 \cdot \text{timeout}_n\}\}$ 
           $\text{unsure}_n \leftarrow t_{\text{now}}$ 
           $\text{freeze}_n \leftarrow t_{\text{now}} + \text{timeout}_n$ 
           $\text{freeze} \leftarrow \max\{\text{freeze}, t_{\text{now}} + t_{\text{powerdown}}\}$ 
         $N_{\text{help}} \leftarrow \emptyset$ 
        // sorgt dafür, dass nicht zu häufig erhöht wird:
         $\text{help\_count} \leftarrow n_{\text{powerup}}$ 
         $\text{help\_repeat} \leftarrow 0$ 
        increase_power()
      else if  $k < k_{\text{min}}$  then
         $\text{help} \leftarrow \text{increase\_k}()$ 
      else if  $k > k_{\text{max}} \wedge t_{\text{now}} > \max\{\text{unsure}_n : n \in N_{\text{bidir}}\} \wedge t_{\text{now}} > \text{freeze}$  then
         $\text{satisfied} \leftarrow \text{decrease\_k}()$ 
    send_beacon( $\text{help}$ ,  $\text{satisfied}$ )
    sleep( $\text{random\_uniform}([1 - t_{\text{jitter}}; 1]) \cdot t_{\text{beacon}}$ )
    // In dieser Zeit: Bearbeitung empfangener Beacons durch recv_beacon()

```

8 Untersuchung der Mechanismen zur Linkerkennung und Sendeleistungsregelung

In diesem Kapitel werden die im vorangegangenen Kapitel 7 eingeführten Mechanismen bewertet. Konkret wird dabei der Frage nachgegangen, unter welchen Umständen welche Varianten des Verfahrens zur verteilten Sendeleistungsregelung in Verbindung mit aktueller Technologie vorteilhaft für den Betrieb drahtloser Netze sind.

Wie bereits in Kapitel 6 festgestellt worden ist, kann Sendeleistungsregelung nur in Netzen mit einer größeren Anzahl von Geräten eine deutliche Kapazitätssteigerung bewirken. Daher ist ein Testbetrieb mit dazugehörigen Messungen nicht mit vertretbarem Aufwand realisierbar. Abgesehen davon hätte dies auch den Nachteil, dass die Bedingungen eines Testlaufs praktisch nicht reproduzierbar sind, was den Vergleich verschiedener Konfigurationen deutlich erschweren würde. Die hier vorgestellten Auswertungen beruhen deshalb auf Simulationsläufen, die Datenübertragungen und Vorgänge, die innerhalb der einzelnen Geräte ablaufen, in einer Modellwelt „durchspielen“. Die verschiedenen Protokolle einschließlich des betrachteten Verfahrens zur verteilten Sendeleistungsregelung werden zu diesem Zweck sehr genau modelliert.

Diese Modellierung wird im folgenden Abschnitt 8.1 genauer erläutert. Abschnitt 8.2 widmet sich dann den Metriken, die in den nachfolgenden Abschnitten der Auswertung der in den Simulationsläufen gewonnenen Daten dienen. Abschnitt 8.3 untersucht den Einfluss verschiedener Parameter auf die Netzkapazität, und Abschnitt 8.4 vergleicht schließlich die einzelnen Varianten der Sendeleistungsregelung miteinander.

8.1 Beschreibung der Simulationsumgebung und der Szenarien

Die Simulationen, deren Ergebnisse im Laufe des Kapitels beschrieben werden, sind mit dem weit verbreiteten Simulationstool ns-2 [FV06] durchgeführt worden. Die einzelnen Parameter sind in Tabelle 8.1 zusammengefasst und werden zusammen mit Modifikationen der Simulationssoftware im Folgenden genauer beschrieben.

Largescale-Fading	Log-Distance-Modell: $n = 2,3$, $d_0 = 300\text{m}$, $\overline{PL}(d_0) = 111\text{dB}$
Smallscale-Fading	Rice-Fading: $K = 4$
RX/CS Threshold	$-91\text{dBm}/-97\text{dBm}$
Sendeleistungen	1mW, 2mW, 3mW, 4mW, 5mW, 6mW, 8mW, 10mW, 15mW, 20mW, 25mW, 30mW, 40mW, 50mW, 75mW, 100mW
Bruttodatenrate	11MBit/s
Data-Link-Layer	IEEE 802.11
ARP	Wird nicht verwendet – Mapping zwischen IP- und MAC-Adressen wird als bekannt vorausgesetzt.
Routing	AODV
Simulationsfläche	quadratisch, $1000\text{m} \times 1000\text{m}$ (in einigen Fällen auch $500\text{m} \times 500\text{m}$ oder $2000\text{m} \times 2000\text{m}$)
Geräte	500 Geräte, gleichverteilt auf der Simulationsfläche
Mobilität	Gauß-Markov-Modell: $\alpha = 0,75$, $\sigma = 0,1\text{m/s}$, $ \mu $ variabel, Update-Intervall 20s

Tabelle 8.1: Simulationsparameter

8.1.1 Bitübertragungsschicht

Die in der Simulationssoftware voreingestellten Werte für die Bitübertragungsschicht beruhen offensichtlich auf sehr alter Hardware und wurden an die Spezifikationen modernerer Geräte angepasst. Als Receive-Threshold wird eine Signalstärke von -91dBm verwendet und als Carrier-Sense-Threshold eine Signalstärke von -97dBm . (Wie in Abschnitt 2.1 beschrieben, hängt der Receive-Threshold von der Kodierung ab; in Abschnitt 8.1.2 wird erklärt, warum nur eine Datenrate, also auch nur eine Modulationstechnik modelliert wird.) Die Sendeleistungen der Geräte können auf 16 verschiedene Werte zwischen 1mW und 100mW eingestellt werden, die in Tabelle 8.1 aufgeführt sind. Dies entspricht dem Wertebereich, der auch von aktueller Hardware (beispielsweise aus der Produktserie „Cisco Aironet 350“ [Cis]) zur Verfügung gestellt wird, wenn auch mit einer geringeren Anzahl möglicher Werte.

Als Modell für das Largescale-Fading dient das Log-Distance-Modell (siehe Abschnitt 2.1.3.1). Lognormal-Shadowing wird nicht berücksichtigt, weil zweifelhaft ist, dass dadurch zusätzliche Erkenntnisse gewonnen werden können, obwohl der dafür benötigte Modellierungsaufwand deutlich höher wäre. Das Lognormal-Shadowing-Modell gibt nämlich die Wahrscheinlichkeit an, mit der eine gewisse Dämpfung der Signalstärke bei vorgegebener Distanz zu erwarten ist. Für zwei konkrete Positionen steht diese Dämpfung aber fest, und die Korrelation zu der Dämpfung, die sich für nur geringfügige Verschiebungen dieser Positionen ergeben, ist relativ hoch. Eine nennenswerte Änderung sollte sich nur dann ergeben, wenn sich die Ausbreitungseigenschaften durch die Positionsverschiebungen deutlich ändern (wenn beispielsweise eine Position um eine Häusercke herum verschoben wird, so dass diese die zuvor offene Sichtlinie zwischen den Positionen verdeckt). Solche Aspekte müssten bei der Modellierung des Lognormal-Shadowings berücksichtigt werden. Daher ist das in ns-2 ent-

haltene Lognormal-Shadowing-Modell für diesen Zweck auch ungeeignet, da dort bei jeder Übertragung eine zufällige Schwankung um den mittleren Pathloss berechnet wird, obwohl es sich ja um ein Largescale-Fading-Modell handelt, das die für gegebene Positionen des Senders und des Empfängers nicht veränderliche mittlere Dämpfung beschreibt.

Der Pathloss-Exponent beträgt 2,3 und entspricht damit einer eher offenen Umgebung mit nur wenigen großen Hindernissen. Die mittlere Dämpfung bei der Referenzdistanz von 300m beträgt 111dB. Mit dem unten aufgeführten Receive-Threshold von -91dBm ergibt sich damit ohne Modellierung des Smallscale-Fadings eine Sendereichweite von 300m für die Sendeleistung von 100mW. Durch den 6dB niedrigeren Carrier-Sense-Threshold ist die Distanz, über die der physikalische Carrier-Sense-Mechanismus das Vorhandensein eines Signals anzeigt, bei dem hier verwendeten Pathloss-Exponenten von 2,3 ungefähr 1,8 mal so hoch wie die Sendereichweite.

Smallscale-Fading wird durch Rice-Fading mit $K = 4 \approx 6\text{dB}$ modelliert (siehe Abschnitt 2.1.3.2). Die dadurch auftretenden Schwankungen der Signalstärke sind als moderat einzustufen; innerhalb von Gebäuden können sie deutlich stärker sein, in offener Umgebung dürften sie dagegen häufig auch wesentlich schwächer sein. In der Realität sind die Linkcharakteristiken im Allgemeinen weniger homogen. Aber auch hier wäre es mit sehr hohem Aufwand verbunden, die Eigenschaften eines Links in Abhängigkeit der Positionen von Sender und Empfänger in einer detailliert modellierten Umgebung zu bestimmen, ohne dass ein qualitativer Einfluss auf die Ergebnisse erwartet werden kann.

Zur Modellierung des Rice-Fading wurde das in [PNS00] beschriebene Modul für den Simulator ns-2 verwendet, das offenbar für Szenarien mit nur einem einzelnen drahtlosen Link konzipiert ist. Es wurde deshalb so angepasst, dass die Signalstärkeschwankungen, die zeitnah entlang verschiedener Links auftreten, unkorreliert sind.

Um den Einfluss des Smallscale-Fadings bewerten zu können, wird ein Teil der Simulationen ohne dessen Modellierung durchgeführt, so dass die Linkzustände binär sind (siehe Abschnitte 2.2.2 und 2.4.1).

8.1.2 Verbindungsschicht

Auf Verbindungsebene wird IEEE 802.11 modelliert (siehe Abschnitt 2.2.1). Weil das in ns-2 enthaltene Modul den Basisstandard in einigen Punkten nur sehr ungenau umsetzt, wurde es neu implementiert. Ebenso so wie mit der ursprünglichen Implementierung wird keine Multirate-Hardware modelliert, in den nachfolgend beschriebenen Simulationen werden daher alle Übertragungen mit 11MBit/s durchgeführt (dies ist die höchste nominelle Datenrate des derzeit sehr weit verbreiteten IEEE 802.11b). Diese vereinfachte Modellierung wurde deshalb beibehalten, weil der im Sinne der Netzkapazität geschickte Einsatz von Hardware zur drahtlosen Datenübertragung, die unterschiedliche Kodierungen und damit unterschiedliche Datenraten unterstützt, eine eigenständige Herausforderung ist, deren Betrachtung über den Rahmen der vorliegenden Arbeit hinausgeht. Dies wurde in Abschnitt 7.2 genauer ausgeführt.

8.1.3 Vermittlungsschicht

Proaktive Routingprotokolle sind für große drahtlose Netze ungeeignet, weil sie mit zu großem Overhead verbunden sind (siehe Abschnitt 2.3.2). Als Routingprotokoll wird daher AODV eingesetzt (siehe Abschnitt 2.3.4). Neben dem Dynamic-Source-Routing (DSR) [JM96] ist es das am weitesten in der IETF-Standardisierung fortgeschrittene reaktive Routingprotokoll für drahtlose Multihop-Netze. Die beiden Protokolle AODV und DSR sind sich in vielerlei Hinsicht so ähnlich, dass sich die Simulationsergebnisse qualitativ kaum unterscheiden dürften. Da die in ns-2 bereits enthaltene Version auf einem älteren Internet-Draft beruht, ist das AODV-Protokoll neu implementiert worden, wobei der aktuelle RFC 3561 [PBRD03] zugrunde gelegt wurde. Die verwendete Implementierung bearbeitet nur solche RREQ-Nachrichten, die über Links empfangen werden, die durch den Beacon-Mechanismus als zuverlässig und bidirektional identifiziert worden sind. Andere RREQ-Nachrichten werden grundsätzlich verworfen.

Das *Address-Resolution-Protokoll* (ARP) [Plu82] ist deaktiviert worden, die Zuordnung von IP-Adressen (Adressen der Ebene 3 des OSI-Referenzmodells [Int94]) zu MAC-Adressen (*Medium-Access-Control*, bezieht sich auf Ebene 2 des OSI-Referenzmodells) wird also als bekannt vorausgesetzt. Diese Entscheidung ist dadurch gerechtfertigt, dass qualitativ keine anderen Ergebnisse zu erwarten wären.

Letztendlich wäre es im Sinne des effizienten Umgangs mit der ohnehin geringen Bandbreite in den betrachteten Szenarien sowieso ratsam, auf den Versand von ARP-Paketen möglichst ganz zu verzichten. Eine denkbare Lösung wäre eine einfache Caching-Strategie: Die als Broadcast-Rahmen versendeten RREQ-Pakete könnten bei allen Empfängern den Eintrag der Verknüpfung zwischen MAC- und IP-Adresse des Senders in der ARP-Tabelle auslösen, so dass für die RREP-Pakete, die entlang der Reverse-Route zurückgesendet werden, keine ARP-Requests mehr nötig sind. Diese RREP-Pakete könnten wiederum die notwendigen ARP-Einträge in Gegenrichtung bewirken, die die für Übertragungen entlang der aufgebauten Route relevanten Informationen enthalten.

Eine solche Caching-Strategie ist nicht konform zu RFC 826 [Plu82], wo vorgeschrieben wird, dass ein neuer Eintrag in der ARP-Tabelle nur durch einen an den Host adressierten ARP-Reply ausgelöst werden darf. Die Motivation dafür ist offenbar, unnötige Einträge in der ARP-Tabelle zu vermeiden. In den hier betrachteten Szenarien könnte man aber durchaus anders zwischen dem Speicherbedarf der ARP-Tabelle und der Menge der versendeten ARP-Pakete abwägen als es bei der Spezifikation des ARP-Protokolls für das klassische Ethernet im Jahr 1982 getan wurde.

8.1.4 Weitere Eigenschaften der beteiligten Geräte

Die untersuchten Szenarien bestehen aus 500 Geräten, die auf einer quadratischen Fläche gleichverteilt positioniert sind. Für die meisten Untersuchungen wird eine Seitenlänge von 1000m verwendet, eine Variation der Gerätedichte erfolgt durch Veränderung dieses Werts

auf 500m bzw. 2000m. Es werden Szenarien mit unterschiedlicher Dynamik untersucht. In den zunächst betrachteten statischen Szenarien bewegen die Geräte sich nicht. In den dynamischen Szenarien werden die Bewegungsvektoren der Geräte gemäß dem in Abschnitt 2.5.2 beschriebenen Gauß-Markov-Mobilitätsmodell im Abstand von 20s neu generiert, wobei die Parameterbelegungen $\alpha = 0,75$ und $\sigma = 0,1\text{m/s}$ verwendet werden. Unterschiedliche Dynamik wird durch Variation von $|\mu|$ erreicht. Falls die Geräte sich auf den Rand der Simulationsfläche bewegen, werden ihr Bewegungsvektor selbst sowie der Vektor μ an der Begrenzung gespiegelt. Wie in Abschnitt 2.5.2 erläutert, sind die Geräte also auch in diesen dynamischen Szenarien zu jeder Zeit gleichverteilt auf der Simulationsfläche positioniert.

In den Szenarien mit aktivierter Sendeleistungsregelung ist jedem Gerät zum Beginn der Simulation zwar eine gemäß einer Gleichverteilung gewählte zufällige Sendeleistung zugeordnet. Um einen möglichst fairen Vergleich zu gewährleisten, wird diese einmalig vorgenommene Zuordnung jedoch für alle Szenarien verwendet.

8.1.5 Verkehrsmuster

Falls die Kapazität der resultierenden Netze untersucht wird, besteht der modellierte Datenverkehr aus 20 Unicast-Strömen zwischen 40 zufällig gewählten, paarweise verschiedenen Geräten. Die Ströme starten im Abstand von etwas über einer Sekunde, so dass nach ca. 30s alle Ströme aktiv sind. Nach weiteren 30s werden alle Ströme zeitgleich beendet. Die Datenrate eines Stroms beträgt 8KBit/s und ist damit ausreichend hoch für die Übertragung von Sprachdaten. Das Paketisierungsintervall beträgt 250ms, wodurch eine für Sprachverbindungen recht hohe Verzögerung zwischen Sprachaufnahme und -wiedergabe entsteht. Dieses Intervall hat jedoch einen starken Einfluss auf die Anzahl der Ströme, die das Netz zur gleichen Zeit transportieren kann: Kleinere Paketisierungsintervalle gehen bei gleicher Datenrate zwar mit geringeren Paketlängen einher, dennoch steigt die Belastung des Netzes durch den mit dem Medienzugang verbundenen Overhead sowie durch die stärkere Konkurrenz zwischen benachbarten Geräten. Deshalb müssen in den hier betrachteten Szenarien höhere Verzögerungen in Kauf genommen werden als in drahtgebundenen Netzen.

8.2 Metriken

8.2.1 Abstand zwischen Zuweisungen von Sendeleistungen

In den folgenden Untersuchungen ist es hilfreich, den Unterschied zwischen zwei Sendeleistungszuweisungen quantifizieren zu können. Damit kann zum einen die Veränderung gemessen werden, die in einem Netz innerhalb einer bestimmten Zeitspanne stattfindet, worauf auch die im nachfolgenden Abschnitt 8.2.2 eingeführte Methode zur Untersuchung von Einschwingzeiten basiert. Zum anderen können die Sendeleistungszuweisungen verglichen werden, die sich aus dem gleichen Startzustand mit unterschiedlichen Algorithmen zur Sen-

deleistungsregelung ergeben.

In den folgenden Auswertungen wird der Fall modelliert, dass allen Geräten die gleiche Menge von n Sendeleistungen $p_1, \dots, p_n$ zur Verfügung steht, wobei $i < j \Leftrightarrow p_i < p_j$. Der Abstand zwischen zwei Sendeleistungszuweisungen $A_1, A_2 : S \rightarrow \{1, \dots, n\}$ wird quantifiziert durch

$$\delta(A_1, A_2) = \sum_{s \in S} |A_1(s) - A_2(s)|,$$

wobei S die Menge aller Geräte ist. Der Wert $\delta(A_1, A_2)$ ist die Anzahl der Sendeleistungsanpassungen auf die nächsthöhere bzw. nächsttiefere Stufe, die notwendig ist, um A_1 in A_2 zu überführen.

Mit dieser Metrik ergibt sich daher auch die Möglichkeit, die *Zielorientierung* von Sendeleistungsanpassungen zu messen: Dazu wird der Quotient aus dem Abstand zwischen Anfangs- und Endtopologie und der Anzahl tatsächlich durchgeführter Sendeleistungsanpassungen gebildet. Im Idealfall werden nur so viele Anpassungen durchgeführt wie nötig, so dass dieser Quotient 1 ist. Je mehr unnötige Anpassungen hinzukommen, also solche, die später rückgängig gemacht werden, desto kleiner wird dieser Quotient.

8.2.2 Einschwingzeiten

Die im vorangegangenen Abschnitt 8.2.1 definierte Abstandsmetrik für Sendeleistungszuweisungen kann auch zur Messung der *Einschwingzeit* dienen. Bezeichnen $[t_0; t_1]$ die Gesamtsimulationszeit und A_t die Sendeleistungszuweisung zum Zeitpunkt $t \in [t_0; t_1]$, so ist der Gesamtabstand gegeben durch $\delta(A_{t_0}, A_{t_1})$, und ein Anteil $x \in [0; 1]$ dieses Gesamtabstands gilt daher zum Zeitpunkt

$$t_{\text{settle}}(x) = \min\{t \in [t_0; t_1] : \delta(A_t, A_{t_1}) \leq (1 - x) \cdot \delta(A_{t_0}, A_{t_1})\}$$

als überwunden, d.h. zu dem Zeitpunkt, an dem der Abstand zur Endzuweisung höchstens $1 - x$ mal so groß ist wie der Gesamtabstand.

In den folgenden Auswertungen wird meist das 90-Perzentil von $t_{\text{settle}}(x)$ berücksichtigt. Dieser Wert gibt an, nach welcher Zeit in 90% der Szenarien ein Anteil x des Gesamtabstands überwunden ist.

8.2.3 Linkpersistenz

Als Maß für die Dynamik der Netztopologie dient die *Linkpersistenz*, die für die folgenden Untersuchungen als Wahrscheinlichkeit dafür definiert wird, dass ein zufällig gewählter Link, der eine Minute vor Ende der Simulationszeit besteht, für die Dauer von einer Minute *persistent* ist. Zur Berechnung der Linkpersistenz wird also die Anzahl jener Links, die die gesamte letzte Minute der Simulationszeit Bestand haben, zur Anzahl aller Links ins Verhältnis gesetzt, die zu Beginn dieses Zeitintervalls bestehen.

Wird kein Smallscale-Fading modelliert, so ist das Ausblenden eines Links aufgrund mehrerer nicht empfangener Beacons in Folge mit hoher Wahrscheinlichkeit darauf zurückzuführen, dass der betroffene Link wegen der Mobilität der Geräte oder wegen einer Sendeleistungsanpassung tatsächlich nicht mehr verfügbar ist.

Wird Smallscale-Fading modelliert, so wird die Linkpersistenz auch davon beeinflusst, wie stark Linkzustände fluktuieren, obwohl die Sendeleistungen unverändert bleiben. In diesem Fall ist die Linkpersistenz lediglich interessant, um die Dynamik der Netztopologie einzuschätzen, wie sie von den Geräten wahrgenommen wird. Es ist aber darauf hinzuweisen, dass eine geringe Linkpersistenz hier keine direkten negativen Auswirkungen hat, da Kommunikationspfade erhalten bleiben, so lange die durchgeführten Übertragungen innerhalb einer vorgegebenen Anzahl von Versuchen erfolgreich sind, unabhängig davon, ob einzelne Links nach der Etablierung eines Pfads von der Linkerkennung ausgeblendet werden. Lediglich für neu zu etablierende Kommunikationspfade werden solche Links nicht mehr berücksichtigt.

8.2.4 Kapazität

Der Durchsatz von Unicast-Verbindungen in drahtlosen Multihop-Netzen wird von verschiedenen Faktoren beeinflusst. Zum einen spielt die Transportkapazität, die in Abschnitt 6.5 auf abstrakterer Ebene untersucht wurde, selbst eine Rolle. Außerdem muss ein Routingprotokoll dafür sorgen, dass die Datenpakete überhaupt übertragen werden können. Dadurch entsteht zusätzlich Kontrollverkehr, der Broadcast-Charakteristiken aufweist.

Wie bereits in Abschnitt 2.2.1 erläutert, ist das Transportprotokoll TCP in drahtlosen Multihop-Netzen, die auf IEEE 802.11 basieren, nicht sehr effektiv. Insbesondere verdeutlicht der ausführliche Vergleich des Durchsatzes verschiedener TCP-Varianten in [XS02], wie kompliziert das Zusammenspiel zwischen TCP und IEEE 802.11 ist. Die Beschränkung des Sendefensters auf wenige Segmente verbessert den Durchsatz von TCP-Verbindungen zwar deutlich, jedoch ist dies ungünstig, um drahtlose Netze zu vergleichen, in denen die Routen je nach verwendeten Sendeleistungen länger und daher mit höherer Verzögerung behaftet oder kürzer und mit geringerer Verzögerung behaftet sind. Die Beschränkung des Sendefensters kann im ersten Fall zu gering sein, so dass die Netzkapazität nicht ausgeschöpft wird, im zweiten Fall kann sie zu hoch sein, um das ungünstige Verhalten von TCP in drahtlosen Netzen zu vermeiden.

Deshalb wird hier die gängige Praxis übernommen, simulative Untersuchungen mit Constant-Bitrate-Verkehr (CBR) durchzuführen, um die Ergebnisse leichter interpretieren zu können und nicht zusätzlich die Wechselwirkungen zwischen TCP und den darunterliegenden Ebenen berücksichtigen zu müssen. In diesem Fall ist es weniger sinnvoll, den Durchsatz selbst zu betrachten, ausschlaggebend ist vielmehr die Paketfehlerrate, die die Anwendung erfährt. In den nachfolgenden Auswertungen wird ein Strom als *erfolgreich* bezeichnet, wenn nach einer Einschwingphase mindestens 80% aller Pakete ihr Ziel erreichen. Die Auswertung der Paketverlustraten der Ströme beginnt nach 30s. Zu diesem Zeitpunkt sind alle Datenströme im hier benutzten Verkehrsmuster (siehe Abschnitt 8.1) gestartet.

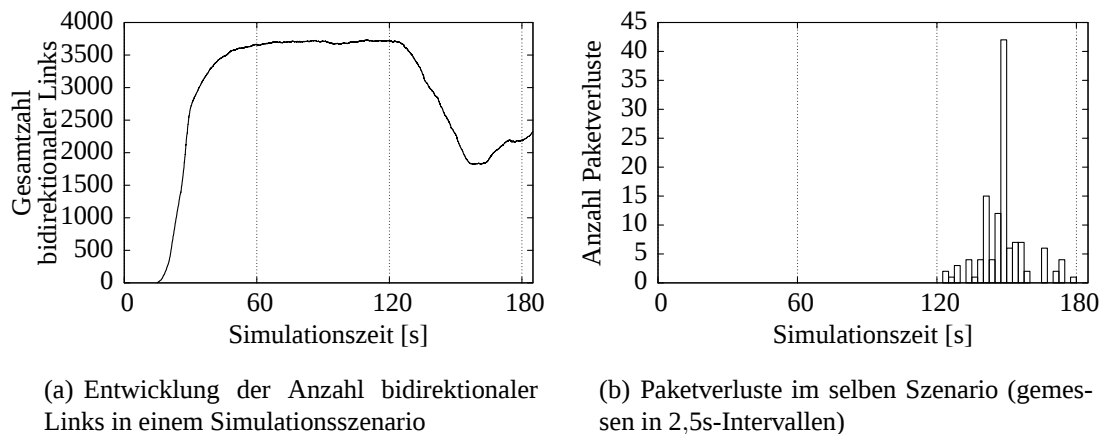


Abbildung 8.1: Anzahl bidirektionaler Links und Paketverluste in einem Beispielszenario im zeitlichen Verlauf

Werden zwei Verfahren bzw. unterschiedliche Parameterbelegungen in einer Reihe von Szenarien miteinander verglichen, so wird der Unterschied zwischen der Anzahl erfolgreicher Ströme dann als signifikant angesehen, wenn das 95%-Konfidenzintervall für den Mittelwert dieses Unterschieds nicht den Wert 0 enthält.

Die genaue Festlegung des Schwellwerts für die maximale Verlustrate eines erfolgreichen Stroms hat auf die Ergebnisse qualitativ keinen großen Einfluss. Insbesondere ist zu beobachten, dass die Paketverlustrate eines Stroms im eingeschwungenen Zustand in den meisten Fällen nahe an 0 oder nahe an 1 liegt. Dies wird im Folgenden genauer erläutert.

Die Beschränkung der Kapazität äußert sich darin, dass aufgrund der starken Konkurrenz um den Medienzugang eine größere Zahl von Übertragungen fehlschlägt, wenn zu viele Datenströme aktiv sind. Erfolgreiche Unicast-Übertragungen werden zwar mehrfach wiederholt, kann ein Paket jedoch auch nach mehreren Versuchen nicht erfolgreich übertragen werden (hier sind es konkret sieben Versuche), so wird der Link als nicht mehr existent betrachtet und eine neue Route wird aufgebaut (im Fall eines Local Repair eine neue Teilroute, siehe Abschnitt 2.3.4). Dadurch entsteht zusätzliche Last, die wiederum zu mehr Paketverlusten führt, und dieser Aufschaukelungskreis mündet schließlich in einem regelrechten „Broadcast-Storm“ [YTCS99]. Einige Ströme können in dieser Situation erfolgreich neue Routen etablieren, und einige Routensuchen schlagen fehl. Das kann daran liegen, dass die lokalen Umgebungen um Quelle oder Senke starker Last ausgesetzt sind, so dass die ungesicherten Broadcast-Rahmen verloren gehen, oder daran, dass zu wenig brauchbare Links vorhanden sind (dies wird im nächsten Absatz genauer erläutert). Fehlgeschlagene Routensuchen werden zwar wiederholt, jedoch finden wegen des exponentiellen Backoff des AODV-Protokolls doch nur relativ wenige Wiederholungen statt, und wenn das Netz an seiner Kapazitätsgrenze arbeitet, können auch diese Wiederholungen erfolglos bleiben. Die Datenpakete dieser Ströme werden dann an ihren Quellen verworfen.

Dies wird auch durch Abbildung 8.1(a) veranschaulicht, die die Anzahl der Links, die durch den Mechanismus zur Linkerkennung als nutzbare bidirektionale Links identifiziert worden sind, in einem Szenario, in dem alle Geräte mit einer einheitlichen, fest vorgegebenen Leistung von 20mW senden, im zeitlichen Verlauf zeigt. (Es handelt sich um eines der in Abschnitt 8.3.1.1 betrachteten Szenarien, es sind $t_{\text{beacon}} = 1\text{s}$, $l_{\text{win}} = 15$, $l_{\text{up}} = 0,9$ und $l_{\text{down}} = 0,7$.) Zunächst dauert es einige Zeit, bis die Anzahl der Links sich auf einen stabilen Wert um etwas unter 4000 einschwingt. Nach 120s beginnt der erste Datenstrom. Die weiteren Datenströme beginnen zeitversetzt, so dass der letzte ca. 25s später startet. Innerhalb der ersten 30s nach Beginn des Datenverkehrs bricht die Anzahl bidirektionaler Links aus dem oben erläuterten Grund bis auf ungefähr die Hälfte ein, und es treten immer mehr Paketverluste auf (siehe Abbildung 8.1(b)). In diesem Szenario können aber für alle Datenströme neue Routen etabliert werden, so dass nach 150s nur vereinzelte Paketverluste auftreten.

8.3 Untersuchung verschiedener Parameterbelegungen

In diesem Abschnitt werden Simulationsergebnisse präsentiert, die die Auswirkungen verschiedener Parameterbelegungen auf Eigenschaften der resultierenden Netze, insbesondere deren Kapazität, zeigen.

Zunächst werden in Abschnitt 8.3.1 Netze betrachtet, in denen die Sendeleistungen der Geräte nicht dynamisch geregelt werden, sondern fest vorgegeben sind. Jedem Gerät wird die gleiche Sendeleistung zugeordnet (analog zu den Common-Range-Topologien aus Kapitel 6). In solchen Netzen werden die Auswirkungen der Linkerkennung (Abschnitt 8.3.1.1), die Auswirkungen variierender Gerätedichte (Abschnitt 8.3.1.2) und die Auswirkungen variierender Dynamik (Abschnitt 8.3.1.3) zunächst entkoppelt von der Sendeleistungsregelung untersucht.

In Abschnitt 8.3.2 werden dann Netze betrachtet, in denen die Standardvariante des in Kapitel 7 vorgestellten Verfahrens zur Sendeleistungsregelung angewendet wird. Abschnitt 8.3.2.1 beschäftigt sich mit der Auswahl eines geeigneten Zielintervalls $[k_{\text{min}}; k_{\text{max}}]$ für die Nachbarnzahlen der Geräte in statischen Szenarien. Abschnitte 8.3.2.2 und 8.3.2.3 untersuchen dann wieder den Einfluss variierender Gerätedichte und Dynamik auf die resultierenden Netze.

8.3.1 Szenarien ohne Sendeleistungsregelung

8.3.1.1 Linkerkennung

Wie in Abschnitt 7.3 beschrieben, wird die Erkennung zuverlässiger Links durch vier Parameter gesteuert. Die Bewertung von Links findet anhand der letzten l_{win} Beacons statt, die ein Gerät versendet hat. Die Zeit zwischen dem Versenden zweier Beacons ist durch das Beacon-Intervall t_{beacon} bestimmt. Sobald der Anteil erfolgreich empfangener Beacons mindestens l_{up}

beträgt, nimmt der Empfänger an, dass ein zuverlässiger unidirektionaler Link vom Sender der Beacons zu ihm selbst besteht. Fällt dieser Wert unter l_{down} , wird diese Annahme wieder verworfen.

Eine sorgfältige Suche nach der optimalen Parameterbelegung ist angesichts der Größe des Parameterraums ein sehr aufwändiges Unterfangen, und eine optimale Wahl ist nicht universell, sondern hängt von verschiedenen Aspekten des konkreten Anwendungsszenarios ab. Hier spielt u.a. eine Rolle, wie stark die durch Smallscale-Fading verursachten Signalstärkeschwankungen sind, wie schnell die Geräte bewegt werden und welche WLAN-Technologie verwendet wird. Da die Erkennung zuverlässiger Links nicht der Schwerpunkt der vorliegenden Arbeit ist, beschränken sich die nachfolgenden Betrachtungen darauf, eine Parameterbelegung zu identifizieren, die in verschiedenen Szenarien annehmbare Ergebnisse liefert.

Ein geeigneter Wert für l_{win} (der Anzahl der Beacons, auf die sich die Bewertung eines Links bezieht) stellt einen Kompromiss zwischen zwei entgegengesetzten Zielen dar: Einerseits ist die Schätzung der Verlustrate eines Links umso zuverlässiger, je größer die dazu betrachtete Stichprobe der Beacon-Übertragungen ist. Andererseits darf die Linkerkennung nicht zu träge gegenüber Veränderungen sein, so dass l_{win} nicht beliebig groß sein sollte. Hier spielt allerdings nicht nur der Parameter l_{win} selbst eine Rolle, wichtig ist vielmehr der Zeitraum, innerhalb dessen l_{win} Beacons versendet werden, also $t_{\text{beacon}} \cdot l_{\text{win}}$. Um auch in dynamischeren Szenarien hinreichend große Werte für l_{win} verwenden zu können, wird $t_{\text{beacon}} = 1\text{s}$ gesetzt, obwohl dieser Wert für weniger dynamische Szenarien wesentlich größer gewählt werden könnte, um mehr Bandbreite für die Übertragung von Nutzdaten zur Verfügung zu haben.

Auch bei der Wahl der Parameter l_{up} und l_{down} sind Kompromisse einzugehen. Unter der Annahme, dass ein Link durch eine gleichbleibende Wahrscheinlichkeit für einen erfolgreichen Beacon-Empfang gekennzeichnet ist, ist die Anzahl der erfolgreich empfangenen unter den letzten l_{win} versendeten Beacons Bernoulli-verteilt und damit zufälligen Schwankungen unterworfen. Der Quotient aus dieser Zahl und l_{win} dient jedoch als Schätzer für die tatsächliche Erfolgswahrscheinlichkeit. Deshalb sollte l_{up} möglichst groß gewählt werden, damit die Wahrscheinlichkeit nicht zu hoch ist, dass über einen Link mit eher schlechter Qualität in einem Zeitfenster dennoch mindestens $\lceil l_{\text{up}} \cdot l_{\text{win}} \rceil$ Beacons empfangen werden. Andererseits dürfen nicht zu wenige Links überhaupt als zuverlässig eingestuft werden, da das Netz sonst gar nicht oder nur unzureichend verbunden ist, was sich negativ auf die Transportkapazität auswirkt (siehe Abschnitt 6.5). Aus ähnlichen Überlegungen heraus darf l_{down} nicht zu gering sein, damit Links minderer Qualität zügig wieder verworfen werden. Ist l_{down} allerdings zu nahe an l_{up} , ist die Linkauswahl nicht sehr stabil, und die Einstufungen als zuverlässig oder unzuverlässig wechseln häufig.

Diese Überlegungen legen eine dynamische Anpassung der Schwellwerte l_{up} und l_{down} nahe, die einerseits möglichst hoch sein sollten, damit die resultierenden Links möglichst gute Übertragungseigenschaften haben, wobei eine gewisse Mindestzahl von Links überhaupt vorhanden sein sollte, damit eine hinreichende Konnektivität des Netzes gegeben ist. Wie in Abschnitt 2.2.2 erläutert, wird in [CE04] deshalb vorgeschlagen, den Schwellwert der Paketverlustrate mit der Anzahl bereits etablierter Links zu erhöhen. Wie bei den in der vorlie-

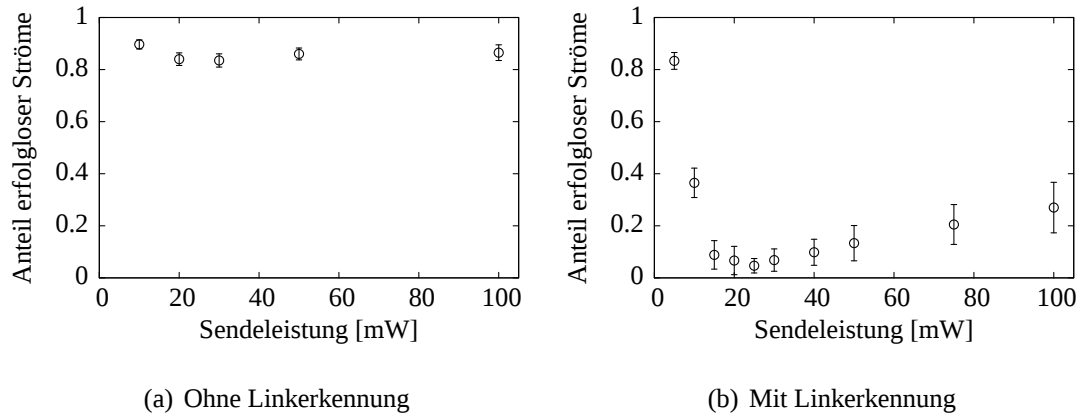


Abbildung 8.2: Erfolg von Datenübertragungen ohne und mit Mechanismus zur Linkerkennung

genden Arbeit untersuchten Mechanismen zur Sendeleistungsregelung geht es dabei um die Formung der Netztopologie zur Kapazitätssteigerung: Falls das Netz über zu wenige Links verfügt (so dass die Konnektivität zu gering ist oder der Routenoverhead zu hoch), so verhilft sowohl eine Verringerung der Schwellwerte zur Linkerkennung als auch eine Erhöhung der Sendeleistungen zu einer höheren Zahl von Links. Es wurde bereits angemerkt, dass eine dynamische Anpassung an die in der Umgebung vorhandene Last hilfreich sein könnte (siehe Abschnitt 7.3). In dem Zusammenhang wurde aber auch festgestellt, dass sich die vorliegende Arbeit auf die Bewertung von Mechanismen zur Sendeleistungsregelung konzentriert. Deshalb ist es sinnvoll, die Linkerkennungsparameter nur mit konstanten Werten zu belegen, um die beobachteten Effekte besser einer Ursache zuordnen zu können.

In dem Fall, dass kein Smallscale-Fading modelliert wird, kann auf das Vorhandensein eines Links geschlossen werden, sobald ein einzelnes Beacon empfangen wird. Erst dann, wenn mehrere Beacons in Folge nicht empfangen werden, muss davon ausgegangen werden, dass ein Link nicht mehr vorhanden ist. Daher werden in diesem Fall die Linkerkennungsparameter mit

$$l_{\text{win}} = 3, \quad l_{\text{up}} = 0,3 \quad \text{und} \quad l_{\text{down}} = 0,3$$

belegt, ein Link wird also wieder ausgeblendet, wenn drei Beacons in Folge nicht empfangen werden.

Aus den im Vorangegangenen beschriebenen Überlegungen heraus werden für Szenarien mit Modellierung von Smallscale-Fading zunächst

$$l_{\text{win}} = 15, \quad l_{\text{up}} = 0,9 \quad \text{und} \quad l_{\text{down}} = 0,7$$

gewählt. Abbildung 8.2 zeigt die 95%-Konfidenzintervalle der Erwartungswerte für den Anteil erfolgloser Datenströme in statischen Szenarien mit Smallscale-Fading. (In Abschnitt 8.2.4 wurde definiert, dass ein Datenstrom *erfolglos* ist, wenn seine Paketverlustrate über 20%

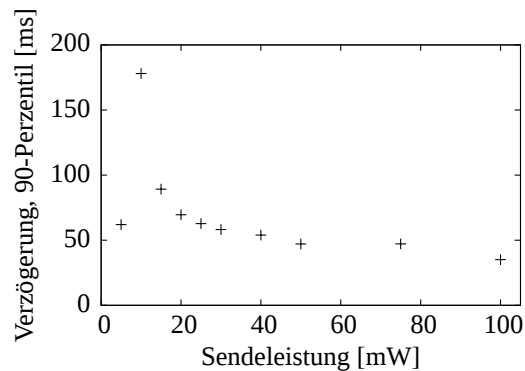


Abbildung 8.3: Verzögerung von Datenpaketen in erfolgreichen Strömen

liegt.) Die Abbildung macht deutlich, wie wichtig ein Mechanismus zur Erkennung zuverlässiger Links unter realistischen Bedingungen ist. In Ermangelung eines solchen Mechanismus ist die Wahrscheinlichkeit hoch, dass das AODV-Protokoll Routen über Links mit hoher Fehlerrate zu etablieren versucht, über die zufälligerweise einzelne RREQ-Pakete übertragen werden konnten. Insbesondere ist zu beachten, dass die Wahrscheinlichkeit für das Fehlschlagen einer Unicast-Übertragung aufgrund des Quittierungsmechanismus von IEEE 802.11 höher ist als die Wahrscheinlichkeit für das Fehlschlagen eines Broadcasts, der zur Übertragung eines RREQ-Pakets verwendet wird. Deshalb kommt es häufig vor, dass eine Route entweder gar nicht aufgebaut werden kann, oder dass ein Link kurze Zeit später vom Routingprotokoll als verloren gegangen betrachtet wird, weil ein Datenpaket nach einer bestimmten Anzahl von Versuchen nicht übertragen werden konnte. An dieser Stelle sei betont, dass diese Problematik nicht auf das AODV-Protokoll beschränkt ist, denn der Grund für die hohen Verlustraten liegt in der vereinfachten Annahme binärer Linkzustände (siehe Abschnitt 2.2.2), die auch den anderen Ad-hoc-Routingprotokollen zugrunde liegt.

Die in Abbildung 8.2(b) dargestellten Ergebnisse decken sich im Übrigen gut mit den Ergebnissen aus Kapitel 6. Auch hier gehen zu geringe ebenso wie zu hohe Sendeleistungen mit höheren Paketverlustraten einher. Bei zu geringen Sendeleistungen liegt das an der geringen Konnektivität, während bei zu hohen Sendeleistungen die starke Konkurrenz um das gemeinsame Übertragungsmedium zunimmt und damit eine geringere Kapazität bewirkt.

Wie Abbildung 8.3 allerdings zeigt, sinkt mit steigender Sendeleistung die Verzögerung der Pakete in erfolgreichen Strömen. Dies mag teilweise auch dadurch zu begründen sein, dass die Gesamtlast geringer ist, weil eine höhere Zahl von Strömen nicht erfolgreich ist und deren Pakete oft schon an ihren Quellen verworfen werden, aber auch beim Vergleich von Szenarien mit Sendeleistungen zwischen 15mW und 30mW, in denen ähnlich viele Ströme erfolgreich sind und ähnlich viele Pakete ihr Ziel erreichen, sind deutliche Unterschiede zu erkennen. Die Anzahl einzelner Übertragungen, die von der Länge der Pfade zwischen Quellen und Senken abhängt, hat also einen wesentlichen Einfluss auf die beobachtete Verzögerung, deren 90-Perzentil jedoch selbst in den Szenarien mit einer einheitlichen Sendeleistung von 15mW noch unter 100ms liegt.

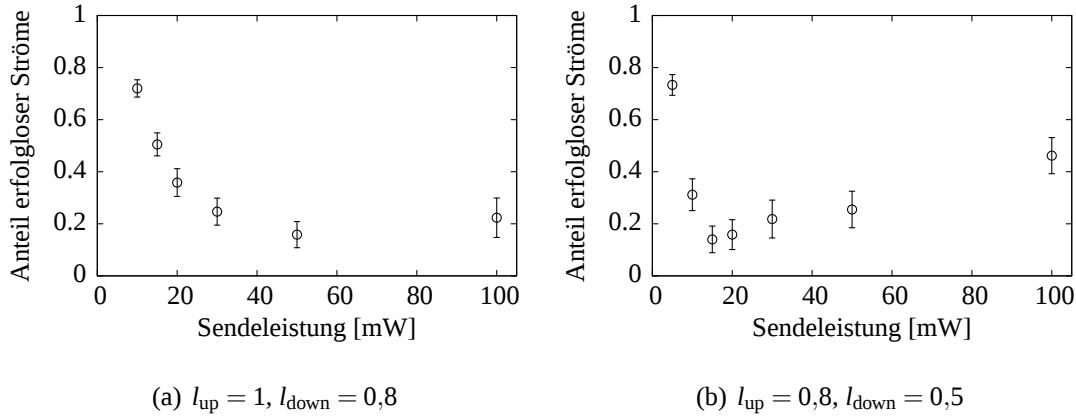


Abbildung 8.4: Erfolg von Datenübertragungen für andere Linkerkennungsparameter

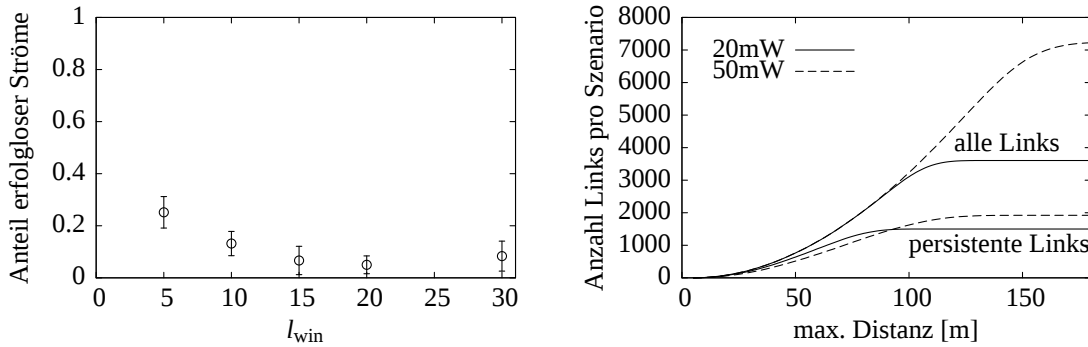


Abbildung 8.5: Erfolg von Datenübertragungen für unterschiedliche l_{win}

Abbildung 8.6: Linkpersistenz in Abhängigkeit von der Distanz

In Abbildung 8.4 ist der Anteil erfolgloser Datenströme für andere Belegungen der Linkerkennungsparameter aufgetragen. Abbildung 8.4(a) zeigt im Vergleich zu Abbildung 8.2(b), dass zu hohe Schwellwerte für die Linkgüte sich insbesondere bei geringeren Sendeleistungen negativ auf die Leistung des Netzes auswirken, da eine zu geringe Anzahl von Links etabliert wird. Abbildung 8.4(b) zeigt weiter, dass zumindest durch die Verringerung der Schwellwerte auf $l_{up} = 0,8$ und $l_{down} = 0,5$ keine besseren Ergebnisse erzielt werden können. Die Ausgangskonfiguration $l_{up} = 0,9$, $l_{down} = 0,7$ ermöglicht insgesamt den geringsten Anteil erfolgloser Datenströme, weshalb diese Werte in allen folgenden Simulationsläufen beibehalten werden.

Abbildung 8.5 zeigt den Anteil erfolgloser Datenströme für unterschiedliche Belegungen von l_{win} , wobei $l_{up} = 0,9$, $l_{down} = 0,7$ sind und die Sendeleistung aller Geräte 20mW beträgt. Offenbar ist die Einschätzung der Linkqualität für $l_{win} \leq 10$ zu ungenau, während sich die Güte der Schätzungen durch $l_{win} \geq 20$ nicht mehr steigern lässt. Für $l_{win} = 30$ ist wieder eine leichte (allerdings nicht signifikante) Erhöhung zu erahnen. Hier liegt die Vermutung nahe, dass zu langsam auf veränderte Bedingungen reagiert wird, in diesem Fall auf das Starten der Datenströme. Durch $l_{win} = 20$ lässt sich zwar die Konfidenz der Ergebnisse ein wenig

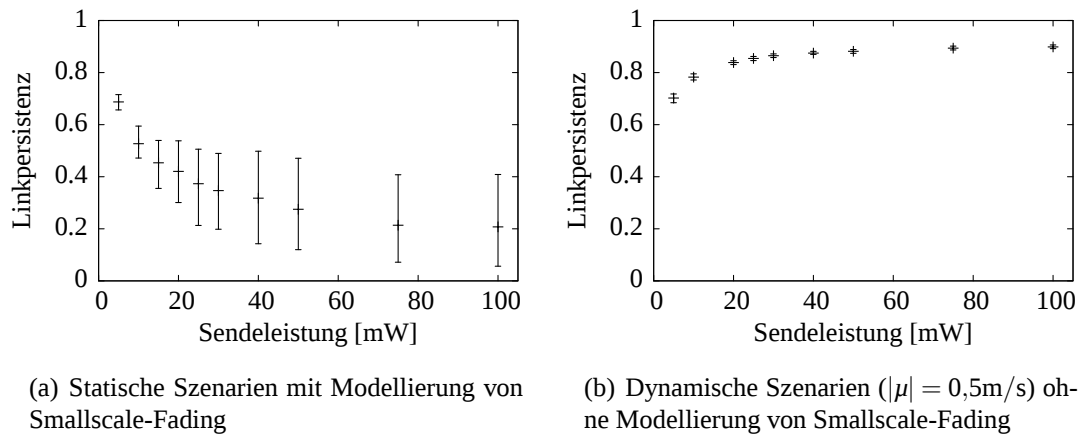


Abbildung 8.7: Linkpersistenz

erhöhen, im Sinne der besseren Anpassungsfähigkeit auch in dynamischen Szenarien wird im Folgenden aber der Einstellung $l_{\text{win}} = 15$ der Vorzug gegeben.

Die aus der Modellierung des Smallscale-Fading resultierende Dynamik wird durch die in Abschnitt 8.2.3 beschriebene Persistenzmetrik verdeutlicht, deren geometrische Mittelwerte zusammen mit 10- und 90-Perzentilen in Abbildung 8.7(a) aufgetragen sind. Warum die Linkpersistenz mit steigenden Sendeleistungen abnimmt, lässt sich anhand von Abbildung 8.6 begründen. Dort ist die Anzahl der Links, die eine Minute vor Ende der Simulationszeit existieren, in Abhängigkeit von der Distanz aufgetragen, die zwischen den beteiligten Geräten liegt. Zusätzlich ist die Zahl der Links aufgetragen, die in der letzten Minute der Simulationszeit persistent sind.

Durch den Unterschied von ca. 4dB zwischen beiden Sendeleistungen 20mW und 50mW und dem Pathloss-Exponenten von 2,3 ergibt sich aus dem Log-Distance-Modell (siehe Abschnitt 2.1.3.1), dass die Distanzen, über die an einem Empfänger eine vorgegebene mittlere Signalstärke erreicht wird, sich für diese Sendeleistungen ungefähr um den Faktor 1,49 unterscheiden. Weil die Anzahl aller Links idealerweise linear von der Größe des Bereichs abhängt, hängt die Quadratwurzel dieser Zahl linear von der Reichweite eines Geräts ab. Setzt man also die Quadratwurzeln der Linkzahlen (ca. 3600 bei 20mW und ca. 7200 bei 50mW, siehe Abbildung 8.6) zueinander ins Verhältnis, ergibt sich ein Wert von ca. 1,42. Die Diskrepanz zu dem aus dem Log-Distance-Modell errechneten Wert von 1,49 ist dadurch zu erklären, dass größere Distanzen bei einer höheren Anzahl von Geräten auch Bereiche abdecken, die außerhalb der Simulationsfläche liegen und daher keine Geräte enthalten, zu denen Links etabliert werden können.

Ein Vergleich der Zahlen persistenter Links in Abhängigkeit der überbrückten Entfernungen für die Sendeleistungen 20mW und 50mW macht deutlich, dass es bei höheren Sendeleistungen wahrscheinlicher ist, dass zwei Beacons ein Gerät zeitlich überlappend erreichen, so dass keines von beiden erfolgreich empfangen werden kann. Dies äußert sich zunächst darin,

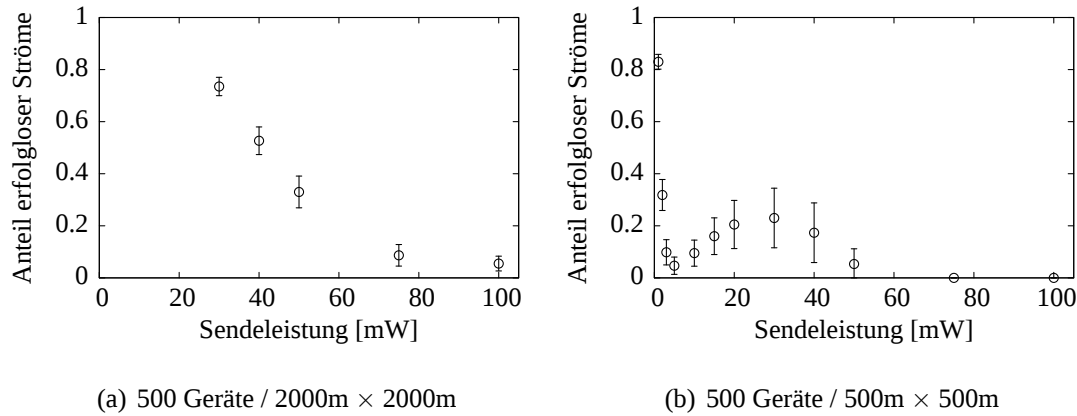


Abbildung 8.8: Erfolg von Datenübertragungen in Szenarien mit anderer Gerätedichte

dass die Anzahl aller Links über Distanzen von bis zu ca. 80m für die Sendeleistungen 20mW und 50mW nahezu gleich ist, dass bei der höheren Sendeleistung jedoch weniger dieser Links persistent sind. Von den Links, die mit der höheren Sendeleistung über größere Distanzen etabliert werden, ist nur ein relativ geringer Anteil persistent.

Als Randbemerkung ist festzuhalten, dass durch eine Vergrößerung des Abstands zwischen den Linkerkennungsparametern l_{up} und l_{down} die Linkpersistenz vergrößert werden könnte. Allerdings sind die hier eingesetzten Werte gerade so gewählt worden, dass sie bei einer geeigneten Wahl der Sendeleistungen zu einer hohen Netzkapazität führen, und es ist bereits dargelegt worden, warum eine Variation dieser Parameter in den folgenden Untersuchungen nicht sinnvoll ist. Dabei ist auch zu bedenken, dass eine geringe Linkpersistenz in den hier betrachteten statischen Szenarien keine unmittelbaren negativen Auswirkungen hat, wie in Abschnitt 8.2.3 erläutert wurde.

8.3.1.2 Gerätedichte

Die Gerätedichte wird durch die Anpassung der Seitenlänge der quadratischen Simulationsfläche variiert. Konkret wird die Seitenlänge auf 2000m erhöht bzw. auf 500m verringert, wodurch sich die Gerätedichte um den Faktor 4 verringert bzw. erhöht. Abbildung 8.8 zeigt den Anteil erfolgloser Datenströme in diesen Szenarien bei Verwendung einheitlicher Sendeleistungen.

Wie zu erwarten ist, sind in den Szenarien mit geringer Gerätedichte höhere Sendeleistungen vorteilhaft (siehe Abbildung 8.8(a)): Der Anteil erfolgloser Datenströme ist mit der maximalen Sendeleistung von 100mW am geringsten.

Ebenfalls erwartungsgemäß sind in den Szenarien mit hoher Gerätedichte bei eher geringen Sendeleistungen (zwischen 3mW und 10mW) nur wenige Ströme erfolglos, während bei höheren Sendeleistungen (im Bereich zwischen 10mW und 30mW) in immer mehr Szenarien

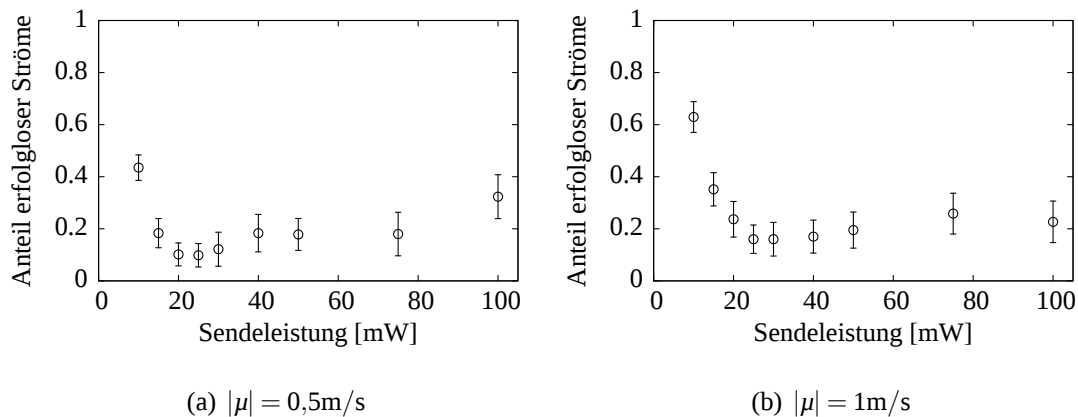


Abbildung 8.9: Erfolg von Datenübertragungen in dynamischen Szenarien

eine Überlastsituation zu beobachten ist. Für noch höhere Sendeleistungen sinkt der Anteil erfolgloser Datenströme allerdings wieder, und bei 75mW und 100mW treten überhaupt keine nennenswerten Paketverluste mehr auf. Diese Ergebnisse zeigen eine gewisse Analogie zu den Ergebnissen aus Abschnitt 6.5, die eine Steigerung der Transportkapazität als Folge erhöhter Sendereichweiten gezeigt haben, falls ohnehin kein Spatial-Reuse mehr möglich ist.

In diesen Szenarien mit hoher Gerätedichte ist die Gesamtkapazität eindeutig höher, wenn größtmögliche Sendeleistungen verwendet werden (und damit auf Spatial-Reuse verzichtet wird), als wenn durch den Gebrauch moderater Sendeleistungen ein höherer Spatial-Reuse angestrebt wird. Dieser Widerspruch zu den Ergebnissen aus Abschnitt 6.5 ist dadurch zu erklären, dass IEEE 802.11 offenbar nicht in der Lage ist, das durch den höheren Spatial-Reuse gegebene Potenzial auszunutzen. Die kurzen Routen und die daraus resultierende geringere Anzahl von Übertragungen bei hohen Sendeleistungen (bei der einheitlichen Sendeleistung von 75mW werden ca. 2/3 aller Pakete über höchstens zwei Zwischenstationen zu ihrer Senke transportiert) sind daher von Vorteil. Mit verbesserten Data-Link-Protokollen ist im Hinblick auf die in Kapitel 6 gewonnenen Erkenntnisse aber nicht auszuschließen, dass auch in solchen Szenarien wie den hier betrachteten geringere Sendeleistungen von Vorteil sind, die einen gewissen Grad an Spatial-Reuse ermöglichen.

8.3.1.3 Dynamik

Wie in Abschnitt 8.1.4 erläutert, erfolgt die Mobilitätsmodellierung gemäß dem in Abschnitt 2.5.2 beschriebenen Gauß-Markov-Modell. Die Dynamik wird dabei durch den Parameter $|\mu|$ festgelegt, also durch den Betrag der (abgesehen von Reflektionen am Rand der Simulationsfläche) unveränderlichen Vektoren, die den einzelnen Geräten zugeordnet sind und als langfristige Mittelwerte der tatsächlichen Bewegungsvektoren dienen.

Abbildung 8.9 zeigt den Anteil erfolgloser Datenströme in dynamischen Szenarien mit $|\mu| =$

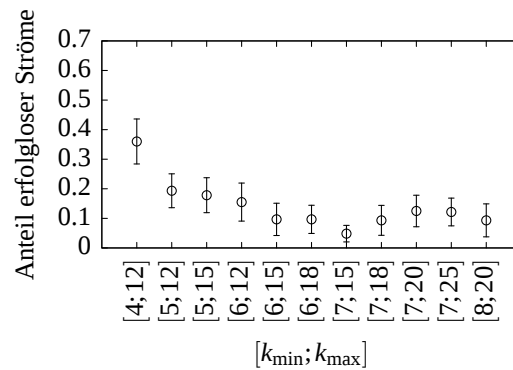


Abbildung 8.10: Erfolg von Datenübertragungen mit verteilter Sendeleistungsregelung

0,5m/s und $|\mu| = 1\text{m/s}$. Während für $|\mu| = 0,2\text{m/s}$ noch keine Unterschiede zur Kapazität in statischen Netzen erkennbar sind, ist im Vergleich zu den in Abbildung 8.2(b) aufgetragenen Anteilen erfolgloser Ströme in statischen Szenarien zu sehen, dass die Netzkapazität mit steigender Mobilität vor allem im Bereich geringerer Sendeleistungen deutlich abnimmt. Die Ursache hierfür ist, dass ein Link zwischen zwei Geräten mit umso höherer Wahrscheinlichkeit innerhalb eines bestimmten Zeitraums mobilitätsbedingt abbricht, je geringer ihre Sendereichweiten sind (siehe auch [GdWFM02]). Damit sinkt auch die Lebensdauer der gewählten Routen, so dass zum einen Datenpakete häufiger verworfen werden und zum anderen die Menge des Signalisierungsverkehrs ansteigt, wodurch weniger Kapazität für den Datenverkehr bleibt. Die durch Erhöhung des Spatial-Reuse erzielbare Steigerung der Netzkapazität fällt daher umso geringer aus, je dynamischer das Netz ist. Abbildung 8.9(b) zeigt, dass in den Szenarien mit $|\mu| = 1\text{m/s}$ schon kein nennenswerter Vorteil gegenüber der Verwendung maximaler Sendeleistungen mehr möglich ist.

Obwohl Geschwindigkeiten von bis zu 1m/s zunächst gering erscheinen mögen, ist die Dynamik der betrachteten Szenarien durchaus beachtlich: Wird als Sendereichweite eines Geräts 100m angenommen, so gehen bei 0,2m/s im Durchschnitt über 6 Links pro Sekunde verloren, bei 0,5m/s über 14 und bei 1m/s fast 28. (Diese Zahlen sind nicht durch Simulationen, sondern nur durch Auswertung der Bewegungsdaten entstanden: Ein Link gilt dann als verloren gegangen, wenn zwei Geräte gegenseitig ihre Sendereichweite verlassen.)

Der beschriebene Zusammenhang zur Lebensdauer der Kommunikationspfade äußert sich auch deutlich in der Linkpersistenz (siehe Abschnitt 8.2.3) in Szenarien, in denen das Smallscale-Fading nicht modelliert wird (Abbildung 8.7(b)). Im Gegensatz zu den statischen Szenarien mit Smallscale-Fading nimmt die Linkpersistenz hier mit höheren Sendeleistungen aus den oben genannten Gründen zu.

8.3.2 Szenarien mit Sendeleistungsregelung

8.3.2.1 Zielintervall der Nachbarzahlen

Abbildung 8.10 zeigt die mittlere Anzahl erfolgloser Datenströme bei Verwendung der Standardvariante des in Kapitel 7 beschriebenen Algorithmus zur Sendeleistungsregelung für mehrere Paare $(k_{\min}, k_{\max})$. Verschiedene Zielbereiche für die Nachbarzahl, die nach den in Abschnitt 6.5 vorgestellten Auswertungen eine sinnvolle Wahl sein könnten, resultieren in einer vergleichsweise geringen Häufigkeit erfolgloser Datenströme. Diese Werte sind für viele der untersuchten Parameterbelegungen ähnlich und sind vergleichbar mit denen, die durch eine geschickte Wahl der einheitlichen Sendeleistung erzielbar sind. In den meisten Fällen sind sie zumindest geringer, als wenn alle Geräte mit maximaler Leistung sendeten (siehe Abbildung 8.2(b)). Diese Tatsache deckt sich wiederum gut mit den Ergebnissen, die in Abschnitt 6.5 präsentiert wurden. Als besonders vorteilhaft stellt sich hier die Parameterbelegung $k_{\min} = 7$ und $k_{\max} = 15$ heraus. Der Anteil erfolgloser Ströme ist mit diesen Parametern nicht signifikant verschieden von denen, die mit einheitlichen Sendeleistungen von 20mW, 25mW oder 30mW auftreten, aber signifikant geringer als mit anderen einheitlichen Sendeleistungen. Das 90-Perzentil der Paketverzögerungen in erfolgreichen Strömen beträgt etwas weniger als 70ms und ist damit nicht höher als in Szenarien mit der einheitlichen Sendeleistung von 20mW (siehe Abbildung 8.3).

Die Linkpersistenz beträgt in diesen Szenarien im Mittel ca. 0,6 und ist damit wesentlich höher als mit festen, einheitlichen Sendeleistungen, die zu einer hinreichenden Konnektivität führen (siehe Abbildung 8.7(a)), obwohl es im gesamten Verlauf der Simulation immer wieder zu Sendeleistungsanpassungen kommt, worauf in Abschnitt 8.4.2 detailliert eingegangen wird. Hier dominiert offensichtlich die Tatsache, dass Beacons mit vergleichsweise geringer Leistung versendet und dadurch nur wenige Störungen verursacht werden. Wegen der hohen Dynamik, der die Netztopologie nur durch das Smallscale-Fading unterliegt, ist die Linkpersistenz in Szenarien mit Smallscale-Fading allerdings nur bedingt dazu geeignet, die zusätzliche Dynamik zu erfassen, die ein Verfahren zur Sendeleistungsregelung erzeugt. Es bleibt aber festzuhalten, dass die von den Geräten wahrgenommene Dynamik bei aktivierter Sendeleistungsregelung sogar eher geringer ist.

8.3.2.2 Gerätedichte

In den Szenarien mit geringerer Gerätedichte ist der Anteil erfolgloser Datenströme mit diesen Parametern ($k_{\min} = 7$ und $k_{\max} = 15$) nicht signifikant verschieden von denen mit einer einheitlichen Sendeleistung von 100mW (siehe Abb. 8.8(a)) und liegen signifikant unter denen, die mit anderen einheitlichen Sendeleistungen auftreten. In den Szenarien mit hoher Gerätedichte ist der Anteil erfolgloser Ströme ähnlich dem im lokalen Optimum im Bereich geringer einheitlicher Sendeleistungen (siehe Abb. 8.8(b)). Dies zeigt, dass das verteilte Verfahren zur Sendeleistungsregelung die Netztopologie in statischen Szenarien so an die Gerätedichte anpasst, dass ein möglichst hoher Spatial-Reuse bei gleichzeitiger Wahrung einer hinreichenden

Konnektivität erreicht wird.

Wie schon für die Szenarien mit einheitlichen Sendeleistungen in Abschnitt 8.3.1.2 festgestellt wurde, ist dies bei hoher Gerätedichte nicht optimal; mit einheitlichen Sendeleistungen von 75mW oder 100mW verringert sich daher die Zahl erfolgloser Ströme signifikant gegenüber der Verwendung der verteilten Sendeleistungsregelung.

8.3.2.3 Dynamik

Wie angesichts der in Abschnitt 8.3.1.3 ausgeführten Tatsachen zu erwarten ist, ist eine deutliche Steigerung der Netzkapazität durch Erhöhung des Spatial-Reuse nur bei moderater Dynamik möglich. Während also in den Szenarien mit $|\mu| = 0,2\text{m/s}$ und $|\mu| = 0,5\text{m/s}$ für keine einheitliche Sendeleistung ein signifikant geringerer Anteil erfolgloser Datenströme beobachtet wird, ist dies in den Szenarien mit $|\mu| = 1\text{m/s}$ für Sendeleistungen zwischen 25mW und 40mW der Fall.

In den Szenarien mit $|\mu| = 0,5\text{m/s}$ führen die vergleichsweise geringen Sendeleistungen bei aktivierter Sendeleistungsregelung bereits zu einer recht niedrigen Linkpersistenz von nur ca. 0,7, wenn kein Smallscale-Fading modelliert wird (zum Vergleich mit Szenarien mit einheitlichen Sendeleistungen siehe Abbildung 8.7(b)). Da die Linkpersistenz in statischen Szenarien ohne Smallscale-Fading nahezu 1 ist (siehe auch Abschnitt 8.4.1), ist der Grund dafür darin zu suchen, dass die Sendeleistungen möglichst gering gehalten werden, um einen hohen Spatial-Reuse zu ermöglichen, so dass die Links nur vergleichsweise kurze Lebensdauern haben. In dynamischen Szenarien mit Smallscale-Fading liegt die Linkpersistenz bei aktivierter Sendeleistungsregelung wieder eher im oberen Bereich der Persistenzwerte, die in Szenarien mit festen einheitlichen Sendeleistungen auftreten, weil die geringen Sendeleistungen bewirken, dass weniger interferenzbedingte Beaconverluste auftreten.

Man könnte argumentieren, dass die Werte $k_{\min} = 7$ und $k_{\max} = 15$, die sich als geeignet für statische Szenarien erwiesen haben, für höhere Dynamik zu gering sind. Um der geringen Linkpersistenz entgegenzuwirken, sollten daher höhere Sendeleistungen verwendet werden. Tatsächlich sind bei $|\mu| = 1\text{m/s}$ mit $k_{\min} = 10$ und $k_{\max} = 18$ weniger Datenströme erfolglos, und deren Anteil ist nicht signifikant höher als in Szenarien mit einheitlichen Sendeleistungen. In Abschnitt 8.3.1.3 wurde aber schon festgestellt, dass diese Dynamik in der vorliegenden Konfiguration bereits so hoch ist, dass geringere einheitliche Sendeleistungen kaum noch einen Vorteil gegenüber maximalen Sendeleistungen bringen. Daher kann auch von einem Verfahren zur Sendeleistungsregelung keine nennenswerte Kapazitätssteigerung erwartet werden.

8.4 Vergleich verschiedener Varianten der verteilten Sendeleistungsregelung

In den folgenden Abschnitten werden die Verhaltensweisen unterschiedlicher Varianten der dezentralen Sendeleistungsregelung untersucht, die in Abschnitt 7.6 beschrieben wurden. Dabei steht die resultierende Kapazität nicht so sehr im Vordergrund wie im vorangegangenen Abschnitt; vielmehr soll ein detaillierter Einblick in das Verhalten der einzelnen Protokollvarianten und die resultierenden Unterschiede gegeben werden. Zu diesem Zweck werden auch Szenarien ohne Smallscale-Fading berücksichtigt. Diese Szenarien lassen aufgrund der hohen Stabilität der Links sehr viel mehr gleichzeitige Ströme zu, so dass auf die Kapazität der Netze nur in Szenarien mit Smallscale-Fading eingegangen wird.

Wie in Abschnitt 8.3.1.1 erläutert, werden zur Linkerkennung die Parameter

$$l_{\text{win}} = 3, \quad l_{\text{up}} = 0,3, \quad l_{\text{down}} = 0,3$$

verwendet, wenn kein Smallscale-Fading modelliert wird, und andernfalls die Parameter

$$l_{\text{win}} = 15, \quad l_{\text{up}} = 0,9, \quad l_{\text{down}} = 0,7.$$

Angesichts der relativ langen Zeitspanne, die verstreichen kann, bis ein Link mit einer Fehlerwahrscheinlichkeit um l_{down} ausgeblendet wird, wird in diesen Szenarien $t_{\text{powerdown}} = 30\text{s}$ gewählt, während in Szenarien ohne Modellierung von Smallscale-Fading $t_{\text{powerdown}} = 5\text{s}$ ausreichend lang ist.

8.4.1 Statische Szenarien ohne Smallscale-Fading

Die in diesem Abschnitt vorgestellten Ergebnisse beziehen sich auf statische Szenarien, in denen kein Smallscale-Fading modelliert wird. Deshalb ist die Signaldämpfung für ein bestimmtes Gerätepaar über die Simulationszeit hinweg konstant und hängt nur von deren Distanz ab.

In Abbildung 8.11 sind die 90-Perzentile von $t_{\text{settle}}(0,95)$ und $t_{\text{settle}}(0,99)$ dargestellt. Diese Werte geben also an, nach welcher Zeit in 90% der untersuchten Szenarien für verschiedene der im Folgenden untersuchten Varianten 95% bzw. 99% des Gesamtabstands zwischen den Sendeleistungszuweisungen zu Beginn und zum Ende der Simulationszeit erreicht sind (siehe Abschnitt 8.2.2). Es zeigt sich, dass mit der Standardvariante, der Max-Degree-Variante und bei Deaktivierung der Linkredundanzsignalisierung die endgültige Topologie schon nach relativ kurzer Zeit annähernd erreicht ist und dass später nur noch wenige Änderungen in größeren Zeitabständen folgen. Die verwendete Simulationszeit von 900s ist in diesen Fällen daher ausreichend hoch. Wird hingegen ausschließlich die Redundanzsignalisierung verwendet, das versuchsweise Reduzieren der Sendeleistungen also deaktiviert, können die Szenarien nach 900s nicht als eingeschwungen bezeichnet werden. Weitere Untersuchungen haben ergeben, dass diese Szenarien selbst nach 3600s kaum eingeschwungen sind, und andererseits sind die

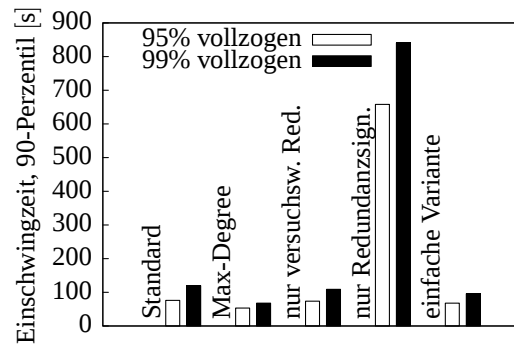


Abbildung 8.11: Einschwingzeiten für verschiedene Varianten des Algorithmus

Aussagen, die im Folgenden über diese Varianten gemacht werden, unabhängig davon, ob die Simulationszeit 900s oder 3600s beträgt.

An dieser Stelle muss angemerkt werden, dass es kein sehr realistisches Szenario ist, in dem 500 Geräte mit zufällig gewählten Sendeleistungen zeitgleich eingeschaltet werden. Sehr viel üblicher wird vermutlich sein, dass in größeren zeitlichen Abständen Geräte zu einem bereits eingeschwungenen Netz hinzugefügt oder aus diesem entfernt werden. Die simulierte Situation ist in dieser Hinsicht also ein Worst-Case-Szenario. Die Einschwingzeiten zeigen dennoch, wie schnell die untersuchten Verfahren in der Lage sind, die Sendeleistungen an veränderte Bedingungen anzupassen.

8.4.1.1 Auswirkung der Kooperation

Die Kooperation ist eine der grundlegenden Innovationen des vorgestellten Algorithmus gegenüber den herkömmlichen Max-Degree-Verfahren: Die Sendeleistung eines Geräts soll nicht nur von dessen eigener Nachbarzahl abhängen, sondern auch die Nachbarzahlen anderer Geräte in der Umgebung berücksichtigen, um ungünstige Situationen, wie sie in Abschnitt 4.2.1.1 beschrieben wurden, zu vermeiden. Hier wird nun untersucht, ob und wie sich diese Kooperation auf die resultierenden Netztopologien auswirkt.

Abbildung 8.12(a) zeigt ein Histogramm über die Anzahl der birektionalen Links, die die Geräte bei Verwendung der Standardvariante des Algorithmus zur Sendeleistungsregelung nach der Simulationszeit von 900s aufgrund des Beacon-Mechanismus feststellen. Fast alle Geräte haben Nachbarzahlen im Bereich $[k_{\min}; k_{\max}]$. Mit der Max-Degree-Variante ist die Verteilung der Nachbarzahlen in diesem Bereich nahezu gleich, so dass das Histogramm für diesen Fall nicht abgebildet ist. Allerdings sind die Häufigkeiten der Nachbarzahlen außerhalb des Zielbereichs deutlich verschieden: In Abbildung 8.12(a) ist zu sehen, dass die Untergrenze $k_{\min}$ nicht unterschritten wird (dies könnte lediglich durch das versuchsweise Reduzieren einer Sendeleistung vorübergehend vorkommen). Dafür ist es erforderlich, dass manche Geräte mehr als $k_{\max}$ Links aufrechterhalten müssen (hier durchschnittlich ca. 0,36% der Geräte). Mit der Max-Degree-Variante wird stattdessen die Obergrenze $k_{\max}$ an keinem Gerät überschrit-

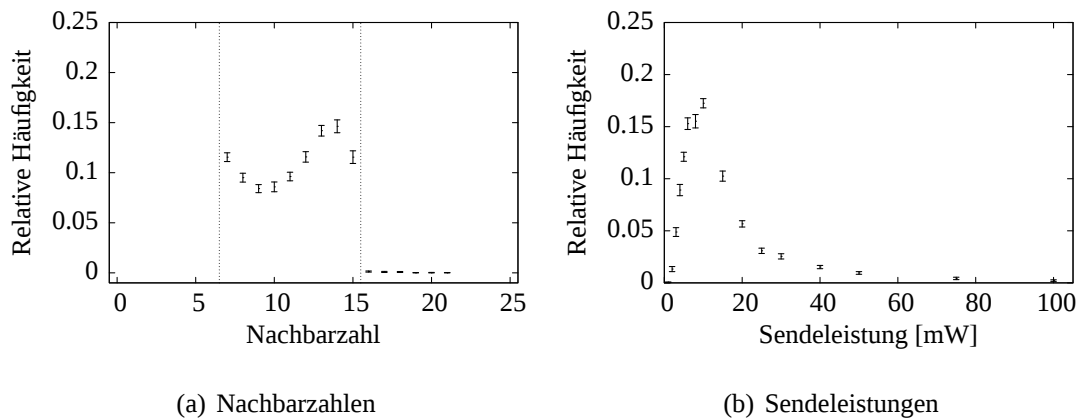


Abbildung 8.12: Nachbarzahlen und Sendeleistungen in statischen Szenarien ohne Smallscale-Fading ($k_{\min} = 7$, $k_{\max} = 15$)

ten, was konsequenterweise jedoch zu Lasten einer geringen Anzahl von Geräten geht, die weniger als $k_{\min}$ bidirektionale Links haben (durchschnittlich sind das ca. 0,65% der Geräte).

Abbildung 8.12(b) zeigt das Histogramm über die von der Standardvariante des Algorithmus eingestellten Sendeleistungen. Wieder unterscheiden sich die Häufigkeiten der Sendeleistungen für die Max-Degree-Variante kaum. In beiden Fällen ist in den untersuchten Szenarien ein erhöhtes Aufkommen von Sendeleistungen im Bereich von ca. 3mW bis 20mW zu beobachten. Es fällt lediglich auf, dass mit der Max-Degree-Variante ein leicht höherer Anteil der Geräte mit maximaler Leistung sendet (im Schnitt 0,9% statt 0,22%). Dies ist größtenteils auf die Geräte zurückzuführen, die weniger als $k_{\min}$ Links haben. Mit dieser Variante erhöht außerdem in fast allen (und zwar in 28 von 30) Szenarien wenigstens ein Gerät seine Sendeleistung von 75mW auf 100mW. Im Schnitt sind zwischen 3 und 4 Geräte (also ca. 0,7% der Geräte) in einem Szenario betroffen. Mit der Standardvariante ist es in den untersuchten Szenarien dagegen sehr selten, dass ein Gerät seine Sendeleistung auf 100mW erhöht (in 27 der 30 Szenarien kommt dies gar nicht vor, und maximal sind 2 Geräte betroffen).

Die beschriebenen Unterschiede betreffen aber nur einige wenige Geräte, und ansonsten sind die erzeugten Topologien nahezu gleich. In den untersuchten Szenarien lohnt sich der Mehraufwand der Kooperation also eher nicht. Etwas deutlicher fallen die Unterschiede dagegen in den Szenarien mit geringerer Gerätedichte aus, in denen die 500 Geräte auf einer quadratischen Fläche mit der Seitenlänge 2000m verteilt sind.

Abbildung 8.13 zeigt die Knotengrade in diesen Szenarien. Mit der Max-Degree-Variante ist der Anteil der Geräte mit weniger als $k_{\min}$ Links hier offenbar höher und liegt durchschnittlich bei immerhin ca. 2,1% (ca. 10 Geräte). Mit der Standardvariante verringert sich dieser Wert auf 0,05%. Hierbei handelt es sich um Geräte, innerhalb deren maximaler Sendereichweite von 300m weniger als $k_{\min}$ andere Geräte platziert sind. Der Anteil der Geräte mit mehr als $k_{\max}$ Links erhöht sich mit der Standardvariante im Schnitt lediglich auf 0,35%. Abbil-

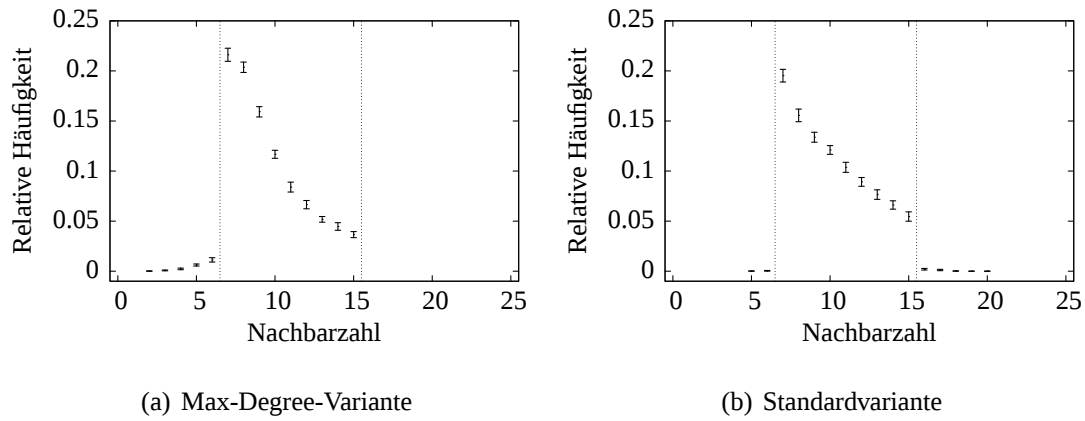


Abbildung 8.13: Nachbarzahlen in statischen Szenarien mit geringer Gerätedichte ohne Smallscale-Fading, $k_{\min} = 7$, $k_{\max} = 15$

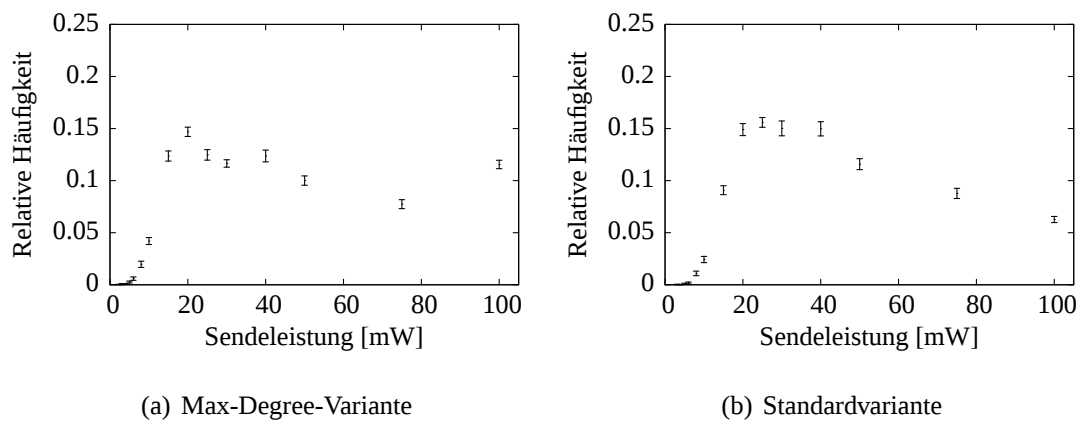


Abbildung 8.14: Sendeleistungen in statischen Szenarien mit geringer Gerätedichte ohne Smallscale-Fading, $k_{\min} = 7$, $k_{\max} = 15$

dung 8.14 zeigt, dass mit der Max-Degree-Variante auch deutlich mehr Geräte die maximale Sendeleistung von 100mW verwenden (ca. 11,6% statt ca. 6,3%). Dies sind also größtenteils Geräte, die zwar die Mindestzahl von $k_{\min}$ Links erreichen können, dabei aber Links über unnötig hohe Entfernungen etablieren müssen.

Der Unterschied zwischen beiden Varianten des Verfahrens manifestiert sich auch in der Inhomogenität der Sendeleistungszuweisungen (siehe Abschnitt 6.3.3). Obwohl die mittlere Zahl bidirektionaler Links mit der Max-Degree-Variante nur etwa 2350 anstatt 2500 mit der Standardvariante beträgt, liegt die mittlere Zahl unidirektionaler Links mit fast 2000 anstatt fast 1600 deutlich höher. Damit ist die Asymmetrie mit der Max-Degree-Variante um ca. 8,2% über derjenigen mit der Standardvariante. In den zuvor betrachteten Szenarien mit höherer Gerätedichte beträgt dieser Unterschied nur ca. 2,9%.

Zuletzt muss noch festgehalten werden, dass die Sendeleistungsanpassungen mit der Standardvariante weniger zielorientiert erfolgen (siehe Abschnitt 8.2.1): Während der Quotient aus dem Abstand der Topologien zum Beginn und zum Ende der Simulationen und der Anzahl tatsächlich durchgeführter Sendeleistungsanpassungen für die Max-Degree-Variante in allen Fällen 1 oder fast 1 ist, liegt er für die Standardvariante wegen des versuchsweisen Reduzierens der Sendeleistungen im Mittel nur bei ca. 0,86.

8.4.1.2 Vergleich der Mechanismen zur Sendeleistungsreduktion

Das vorgestellte Protokoll sieht in der Standardvariante zwei Mechanismen zur Reduzierung der Sendeleistung vor: Geräte mit mehr als $k_{\max}$ Links reduzieren ihre Sendeleistung, wenn ihre Links alle von den Nachbarn als redundant eingestuft werden (siehe Abschnitt 7.5.2), und falls nicht redundante Links davon betroffen sind, reduzieren sie ihre Sendeleistung versuchsweise mit exponentiellem Backoff (siehe Abschnitt 7.5.3). Hier wird nun der Frage nachgegangen, wie diese Mechanismen jeweils für sich betrachtet arbeiten.

Wird nur das versuchsweise Reduzieren als Mechanismus zur Senkung der Sendeleistung verwendet, die Signalisierung redundanter Links also deaktiviert, so ähneln die resultierenden Topologien sehr stark denen der Standardvariante. Dies kann durch die Abstandsmetrik bestätigt werden (siehe Abschnitt 8.2.1): Der durchschnittliche Abstand über alle 30 Szenarien liegt bei nur ca. 12, der Maximalwert beträgt 31.

Abbildung 8.15(a) zeigt ein Histogramm über die Nachbarzahlen der Geräte, wobei das versuchsweise Reduzieren von Sendeleistungen deaktiviert ist und nur die Redundanzsignalisierung zur Sendeleistungsreduktion verwendet wird. Erwartungsgemäß tritt hier häufiger der Fall ein, dass Geräte mit zu vielen Links ihre Sendeleistung reduzieren könnten, ohne dass nicht redundante Links dadurch abbrechen, wie die hohe Zahl von Geräten mit mehr als $k_{\max}$ Nachbarn (fast 30%) im Vergleich zu Abbildung 8.12(a) verdeutlicht.

Abbildung 8.15(b) zeigt schließlich die Nachbarzahlen, die sich durch die Verwendung der einfachen Variante des Verfahrens ergeben, die in [GdWMJ03], [GdWMJ05] vorgestellt wurde und die die Sendeleistung eines Geräts nur dann verringert, wenn dieses selbst sowie alle

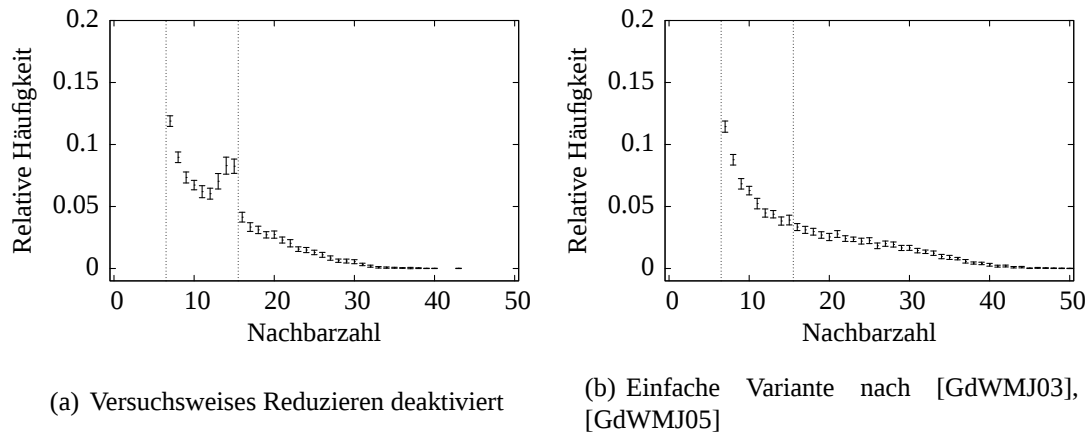


Abbildung 8.15: Sendeleistungen ohne versuchsweises Reduzieren in statischen Szenarien, $k_{\min} = 7$, $k_{\max} = 15$

seine Nachbarn mehr als $k_{\max}$ Links haben. Es ist deutlich zu erkennen, dass diese Variante viel zu häufig nicht in der Lage ist, durch eine Verringerung von Sendeleistungen Nachbarzahlen unter oder auch nur nahe an $k_{\max}$ zu erreichen: Im Mittel haben ca. 45% der Geräte mehr als $k_{\max}$ Nachbarn.

8.4.2 Statische Szenarien mit Smallscale-Fading

Im folgenden Abschnitt wird eine Reihe von Simulationen beschrieben, denen die gleichen statischen Szenarien zugrunde liegen wie im vorangegangenen Abschnitt, die aber unter Berücksichtigung von Smallscale-Fading-Effekten durchgeführt wurde.

8.4.2.1 Einschwingzeiten

Durch die Modellierung des Smallscale-Fadings ist es (im Gegensatz zu den im vorangegangenen Abschnitt untersuchten Szenarien) sehr viel schwieriger, Aussagen über die Einschwingzeit der Szenarien zu treffen. Aufgrund von spontanen Änderungen der Linkzustände kommt es nämlich während der gesamten Simulationszeit zu Anpassungen der Sendeleistungen, wodurch die Netztopologie immer wieder neue stabile Zustände erreicht. Daher müsste die langfristige Entwicklung der Netztopologie betrachtet werden, jedoch wären dazu extrem rechenaufwändige Simulationsläufe notwendig gewesen, deren zu erwartender Erkenntnisgewinn in keinem sinnvollen Verhältnis zum Aufwand steht.

Abbildung 8.16(a) stellt den Verlauf des 90-Perzentils von $t_{\text{settle}}(x)$ über $x \in [0; 1]$ dar (siehe Abschnitt 8.2.1). Aus der Abbildung ist ersichtlich, dass bei der Simulationszeit von 1800s zunächst sehr starke Änderungen an der Netztopologie vorgenommen werden, dass mit bei-

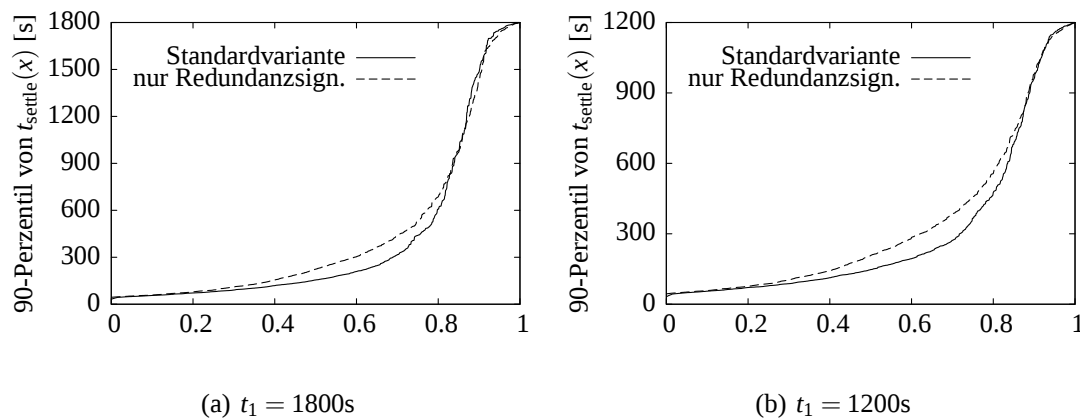


Abbildung 8.16: Verlauf der Topologieentwicklung in Szenarien mit Smallscale-Fading

den Varianten aber nach 900s über 80% des Gesamtabstands überwunden sind. Danach ist der Verlauf der Topologieentwicklung nur noch durch geringfügige Änderungen gekennzeichnet, und auch bei genauerer Betrachtung ist nicht erkennbar, dass die Änderungsfrequenz mit fortschreitender Simulationszeit weiter abnähme, vielmehr schwankt diese nur um relativ geringe Werte herum. Erst in den letzten 200s ist wieder eine schnellere Entwicklung hin zur Zieltopologie zu sehen.

Wie Abbildung 8.16(b) zeigt, tritt dieser Anstieg zum Ende der Simulation aber auch dann auf, wenn ein früherer Endzeitpunkt gewählt wird, in diesem Fall 1200s. Es sei hier betont, dass es sich um die selben Simulationsläufe handelt, die lediglich mit einem früheren Endzeitpunkt ausgewertet wurden. Die Erklärung dieses Anstiegs ist in der Tatsache zu finden, dass die absolute Anzahl tatsächlich durchgeführter Sendeleistungsanpassungen pro Zeiteinheit sich zum Ende der Simulationen nicht erhöht, dass unter den Anpassungen gegen Ende der Simulationen aber ein höherer Anteil zum Erreichen des Endzustands notwendig ist. Dieser Endzustand des sich fortwährend und wenig zielgerichtet ändernden Systems zu einem willkürlich gewählten Endzeitpunkt ist letztlich aber ebenfalls willkürlich. (Die Zielorientierung, die gemäß Abschnitt 8.2.1 als der Quotient zwischen der Anzahl notwendiger und der Anzahl tatsächlich durchgeführter Sendeleistungsanpassungen gemessen wird, liegt für beide Varianten im Mittel unter 0,5.)

Mit diesen Überlegungen lässt sich aus Abbildungen 8.16(a) und 8.16(b) schließen, dass die wesentlichen Änderungen nach den ersten 900s bereits vollzogen sind. Diese Simulationszeit von 900s liegt daher den im Folgenden vorgestellten Simulationsergebnissen zugrunde.

8.4.2.2 Auswirkung der Kooperation

Abbildung 8.17 zeigt Histogramme für die Anzahl bidirektionaler Links, die die Geräte anhand des Linkerkennungs-Mechanismus nach der Simulationszeit von 900s feststellen, für die

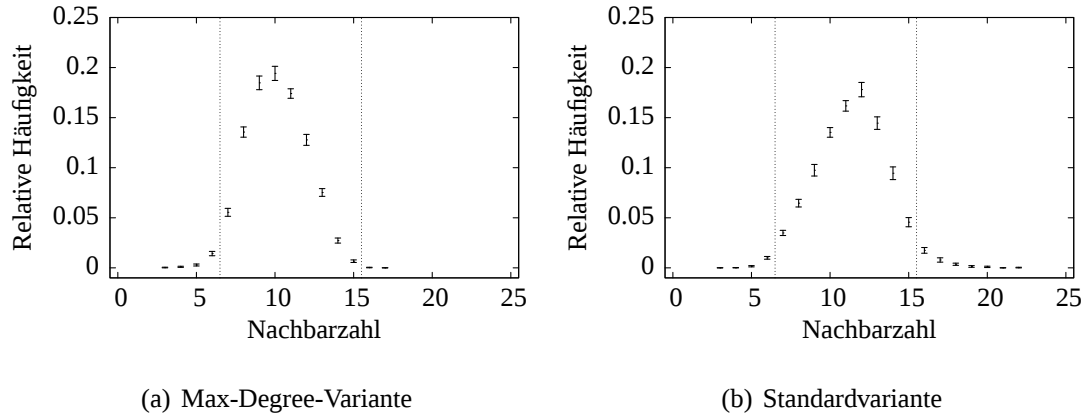


Abbildung 8.17: Knotengrade in statischen Szenarien mit Smallscale-Fading, $k_{\min} = 7$, $k_{\max} = 15$

Standardvariante und die Max-Degree-Variante. Die Histogramme haben ein deutlich anderes Aussehen als ohne die Modellierung von Smallscale-Fading (siehe Abbildung 8.12(a)).

Es ist zu erkennen, dass die Nachbarzahlen glockenförmig im Bereich $[k_{\min}; k_{\max}]$ verteilt sind. Dies ist dadurch zu erklären, dass die Nachbarzahl eines Geräts auch während einer Zeitspanne, in der die Sendeleistungen unverändert bleiben, wegen des Smallscale-Fadings um einen bestimmten Wert herum schwankt. Je näher dieser Wert an $k_{\min}$ oder $k_{\max}$ liegt, desto wahrscheinlicher ist es, dass durch das sporadische Hinzukommen bzw. Wegfallen von Links der Bereich $[k_{\min}; k_{\max}]$ verlassen wird. Dies löst entsprechende Sendeleistungsanpassungen aus, die eine Veränderung des besagten Mittelwerts bewirken.

Mit der Standardvariante sind die Nachbarzahlen etwas höher als mit der Max-Degree-Variante, und durchschnittlich 3,2% der Geräte haben mehr als $k_{\max}$ Links. Dies ist dadurch zu erklären, dass aus dem oben dargelegten Grund zwar sowohl Werte nahe $k_{\min}$ als auch nahe $k_{\max}$ auf längere Dauer instabil sind, dass letztere jedoch weniger stark von der beschriebenen Instabilität betroffen sind, weil der Algorithmus eher vorsichtig bei der Verringerung von Sendeleistungen vorgeht.

Dennoch ist auch hier wieder zu beobachten, dass mit der Max-Degree-Variante deutlich mehr Geräte mit maximaler Leistung senden (ca. 3,4% statt 1,15%). Daraus kann geschlossen werden, dass die Sendeleistungszuweisungen mit dieser Variante wieder inhomogener sind (da wegen des Smallscale-Fadings keine eindeutige Sendereichweite angegeben werden kann, ist die in Abschnitt 6.3.3 vorgestellte Asymmetrie-Metrik hier nicht anwendbar.) Dies ist wahrscheinlich der Grund dafür, dass mit der Max-Degree-Variante signifikant mehr Ströme erfolglos sind als mit der Standardvariante (im Mittel ca. 4 Ströme statt ca. 1 von insgesamt 20). Dies hängt aber auch damit zusammen, dass die Max-Degree-Variante bei gleichem Zielintervall generell eine geringere Anzahl von Links etabliert als die Standardvariante. Durch die Erhöhung der Nachbarzahlparameter auf $k_{\min} = 10$ und $k_{\max} = 18$ verringert sich der An-

teil erfolgloser Ströme auf ca. 2,5, durch eine weitere Erhöhung $k_{\max} = 25$ gar auf 2. Die Unterschiede zur Standardvariante bleiben aber signifikant, und der beobachtete Nachteil der entstehenden Topologien, die hohe Zahl von Geräten, die mit maximaler Leistung senden, wird noch verstärkt.

8.4.2.3 Vergleich der Mechanismen zur Sendeleistungsreduktion

Interessant ist weiterhin die Beobachtung, dass auch der Mechanismus zur Redundanzsignalisierung alleine (also ohne das versuchsweise Reduzieren der Sendeleistungen) in diesen Szenarien zu guten Ergebnissen führt, während ohne Modellierung des Smallscale-Fadings ja eine hohe Zahl von Geräten mehr als $k_{\max}$ Links hatte (siehe Abbildung 8.15(a)). Die Verteilung der Nachbarzahlen nach 900s ist in diesen Szenarien sehr ähnlich zu der in Abbildung 8.17(b) dargestellten, mit einer nur leichten Tendenz zu höheren Werten (durchschnittlich haben 4,4% statt 3,2% der Geräte mehr als $k_{\max}$ Links).

Die stärkere Fluktuation der Linkzustände ist für die Variante ohne versuchsweises Reduzieren der Sendeleistungen hilfreich, weil es wie oben erläutert eher unwahrscheinlich ist, dass ein Gerät über einen längeren Zeitraum hinweg eine Nachbarzahl von genau $k_{\min}$ hat. Gerade solche Geräte hindern ihre Nachbarn nämlich langfristig am Reduzieren der Sendeleistung. Außerdem können nicht-redundante Links kurzzeitig wegfallen, was Geräten mit mehr als $k_{\max}$ Links eine Reduktion der Sendeleistung erlaubt.

Abbildung 8.16(a) zeigt, dass die Konvergenz ohne das versuchsweise Reduzieren der Sendeleistungen ein wenig langsamer verläuft, dafür lässt sie aber eine etwas größere Zielorientierung erkennen: Innerhalb der Simulationszeit von 900s werden im Mittel fast 3000 Sendeleistungsanpassungen durchgeführt, während es mit der Standardvariante im Mittel dagegen schon über 3500 sind. Die Mindestanzahl benötigter Sendeleistungsanpassungen (d.h. der Abstand zwischen Sendeleistungszuweisungen zum Beginn und zum Ende einer Simulation) beträgt in beiden Fällen im Mittel ca. 1900.

Dieser Unterschied zwischen den Varianten macht sich in der Linkpersistenz allerdings nicht bemerkbar, da die Linkpersistenz während einer Zeitspanne am Ende der Simulation gemessen wird, in der die Sendeleistungen also relativ stabil sind und die Timeout-Werte für das versuchsweise Reduzieren der Sendeleistungen ausreichend hoch. Die Linkpersistenz beträgt in beiden Fällen im Mittel ca. 0,6.

Die Ähnlichkeit der aus den verschiedenen Varianten resultierenden Topologien äußert sich auch in der Netzkapazität: Weder die Deaktivierung des versuchsweisen Reduzierens noch die Deaktivierung der Redundanzsignalisierung ändern die Häufigkeit erfolgloser Datenströme signifikant (die im Schnitt bei ca. einem von 20 Strömen liegt).

Wie die Auswahl redundanter Links vollzogen wird (ob mit gleicher Wahrscheinlichkeit für alle Nachbarn oder mit höherer Wahrscheinlichkeit für Nachbarn mit mehr Links) spielt in diesen Szenarien keine erkennbare Rolle.

Schließlich ist die einfache Variante des Algorithmus gemäß [GdWMJ03], [GdWMJ05] auch

in diesen Szenarien nicht in der Lage, den Anteil der Geräte mit mehr als $k_{\max}$ Nachbarn in einem akzeptablen Rahmen zu halten: Genau wie ohne Modellierung des Smallscale-Fadings haben im Schnitt wieder ca. 45% der Geräte zu viele Links. Dadurch verringert sich die Kapazität im Vergleich zur Standardvariante signifikant (im Mittel sind fast drei von 20 Strömen erfolglos).

8.4.3 Dynamische Szenarien ohne Smallscale-Fading

Nachdem in den vergangenen Abschnitten das Verhalten verschiedener Varianten des verteilten Algorithmus zur Sendeleistungsregelung in statischen Szenarien untersucht worden ist, wird in den nächsten beiden Abschnitten deren Verhalten in dynamischen Szenarien untersucht. Zunächst werden, wie in Abschnitt 8.4.1, die Effekte des Smallscale-Fadings ausgeblendet.

Schon die im letzten Abschnitt betrachteten statischen Szenarien hatten durch die Modellierung des Smallscale-Fadings eine gewisse Dynamik, die die Betrachtung von Einschwingzeiten schwierig gemacht hat, weil sich die Netztopologie laufend geändert hat. Weil die Geräte aber nicht bewegt wurden, war die Menge der Sendeleistungszuweisungen, die einen relativ stabilen Zustand des Netzes bewirken, unveränderlich. Die hier betrachteten Szenarien sind jedoch dynamisch in dem Sinn, dass die Geräte laufend ihre Positionen ändern, und eine Sendeleistungszuweisung, die während einer bestimmten Zeitspanne einen relativ stabilen Netzzustand bewirkt, kann in einer anderen Zeitspanne keinen langen Bestand haben. Deshalb müsste die Einschwingzeit eines Szenarios mit mobilen Geräten anders definiert werden. Für die Betrachtungen in dieser Arbeit reicht jedoch die Überlegung, dass eine Zeitspanne von 900s, die ein Vielfaches der kurzen Einschwingzeiten in statischen Szenarien ohne Smallscale-Fading ist (siehe Abbildung 8.11)), zum Erreichen eines eingeschwungenen Zustands in Szenarien mit mobilen Geräten ebenfalls ausreichen muss.

Ähnlich wie in statischen Szenarien mit Smallscale-Fading (siehe Abschnitt 8.4.2) sind die Nachbarzahlen hier eher glockenförmig verteilt (siehe Abbildung 8.18). Die Ursache ist prinzipiell die gleiche: Je näher die Nachbarzahl eines Geräts am Rand des Zielintervalls ist, desto wahrscheinlicher ist es, dass dieses Intervall innerhalb eines festen Zeitraums verlassen wird, was entsprechende Sendeleistungsanpassungen bewirkt. Im einen Fall liegt dies an Fluktuationen von Linkzuständen, die durch das Smallscale-Fading bedingt sind, im anderen Fall an mobilitätsbedingten Topologieänderungen.

8.4.3.1 Vergleich der Mechanismen zur Sendeleistungsreduktion

Eine weitere Analogie zum vorangegangenen Abschnitt ist die Tatsache, dass auch die Variante ohne das versuchsweise Reduzieren der Sendeleistungen relativ effektiv ist. Abbildung 8.18(a) zeigt die Nachbarzahlen in den Szenarien mit $|\mu| = 0,5\text{ms/s}$; hier haben im Schnitt 8,2% der Geräte mehr als $k_{\max}$ Links (anstatt 6,5% bei der Standardvariante). Die Verteilung

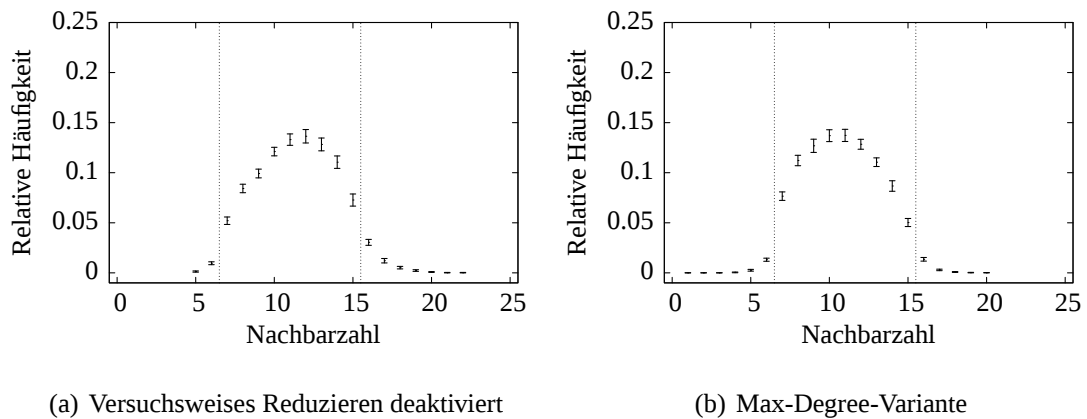


Abbildung 8.18: Knotengrade in dynamischen Szenarien ohne Smallscale-Fading, $|\mu| = 0,5\text{m/s}$, $k_{\min} = 7$, $k_{\max} = 15$

der Nachbarzahlen für die Standardvariante verläuft insbesondere für Werte bis $k_{\max}$ sehr ähnlich.

Abbildung 8.19(a) zeigt die Verteilung der Sendeleistungen für den Fall, dass das versuchsweise Reduzieren der Sendeleistungen deaktiviert ist (die Verteilung für die Standardvariante ist auch hier nahezu gleich). Gegenüber statischen Szenarien (siehe Abbildung 8.12(b)) sind die Sendeleistungen deutlich homogener: Während mehr Geräte mit Leistungen zwischen 5mW und 8mW senden, ist der Anteil geringer Sendeleistungen (3mW und weniger) und hoher Sendeleistungen (20mW und mehr) gesunken.

Die Situation ist hier offenbar anders als in statischen Szenarien, in denen die Sendeleistungen nur so lange angepasst werden, bis ein stabiler Zustand erreicht ist. Ein solcher Zustand zeichnet sich dadurch aus, dass kein Gerät weniger als $k_{\min}$ Links hat (sofern dies überhaupt möglich ist) und dass die Geräte mit mehr als $k_{\max}$ Links ihre Sendeleistung nicht reduzieren können, ohne dass dadurch Links verloren gingen, auf die ihre Nachbarn angewiesen sind. Das schließt aber nicht aus, dass Geräte mit einer Nachbarzahl im vorgegebenen Zielbereich dennoch mit relativ hoher Leistung senden. Durch die Dynamik des Netzes kommen auch solche Geräte in den hier betrachteten Szenarien mit großer Wahrscheinlichkeit in Bereiche mit höherer Gerätedichte und reduzieren ihre Sendeleistungen daraufhin. Analog dazu kommen Geräte, die im statischen Fall bei eher geringen Sendeleistungen verweilen, in weniger dichte Bereiche des Netzes, woraufhin ihre Sendeleistungen erhöht werden müssen. Auf diese Weise pendeln sich die Sendeleistungen auf den Stufen ein, die der Gesamtdichte des Netzes angemessen sind. Dies äußert sich auch in der verringerten Asymmetrie der erzeugten Topologien: Während die mittlere Anzahl bidirektionaler Links mit ca. 2800 im Vergleich zu den in Abschnitt 8.4.1 betrachteten statischen Szenarien fast unverändert ist, verringert sich die mittlere Anzahl unidirektionaler Links sehr deutlich von ca. 2500 auf ca. 1300.

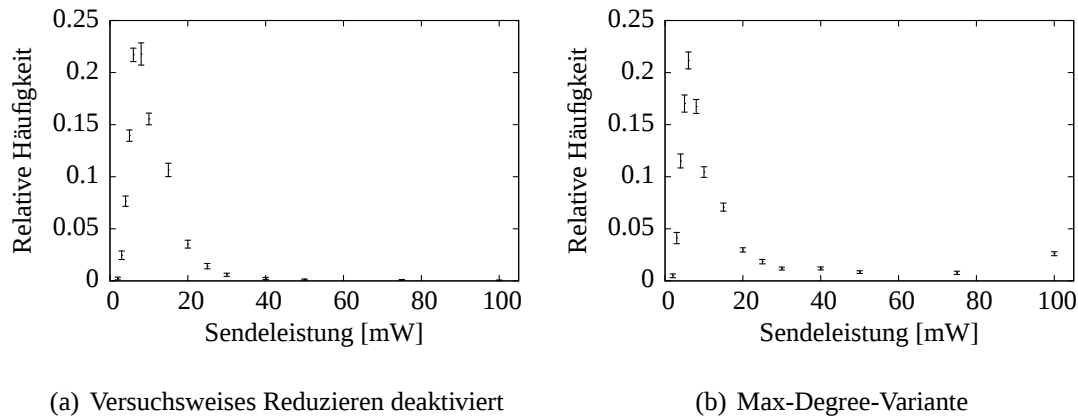


Abbildung 8.19: Sendeleistungen in dynamischen Szenarien ohne Smallscale-Fading, $|\mu| = 0,5\text{m/s}$, $k_{\min} = 7$, $k_{\max} = 15$

8.4.3.2 Auswirkung der Kooperation

Für die Max-Degree-Variante ergibt sich ein anderes Bild, wie die in Abbildung 8.19(b) dargestellte Verteilung der Sendeleistungen zeigt. Ähnlich wie für die Standardvariante sinkt gegenüber statischen Szenarien der Anteil geringerer Sendeleistungen (3mW und weniger) und höherer Sendeleistungen (zwischen 10mW und 50mW). Allerdings ist ein deutlicher Anstieg des Anteils der Geräte zu beobachten, die die maximale Sendeleistung von 100mW verwenden, und dies obwohl die Verteilung der Nachbarzahlen eine leichte Tendenz zu geringeren Werten aufweist als mit der Standardvariante (siehe Abbildung 8.18). Daher ist mit der Max-Degree-Variante (anders als mit der Standardvariante) auch keine nennenswerte Änderung der Asymmetrie gegenüber statischen Szenarien zu beobachten. Sowohl die mittlere Zahl bidirektionaler Links als auch die mittlere Zahl der unidirektionalen Links sind im wesentlichen unverändert (beide betragen ca. 2700). In den hier betrachteten dynamischen Szenarien ist die Anzahl unidirektionaler Links im Mittel also ungefähr doppelt so hoch wie mit der Standardvariante, obwohl die Zahl bidirektionaler Links im Mittel fast gleich ist.

Dieses Verhalten ist eine Folge davon, dass durch die Dynamik des Netzes mehr Geräte im Verlauf einer Simulation irgendwann in die Situation geraten, nicht genügend Links aufbauen zu können, wodurch sie veranlasst werden, mit maximaler Leistung zu senden: Für die Max-Degree-Variante fällt im Vergleich zu den statischen Szenarien auf, dass die Geräte deutlich häufiger ihre Sendeleistung von 75mW auf den Maximalwert 100mW erhöhen müssen. Im Durchschnitt über alle 30 Szenarien sind davon ca. 40 Geräte (also ca. 8% aller Geräte) betroffen, und in allen Szenarien sind es über 20 Geräte. Mit der Standardvariante ist dies dagegen nur in einem Szenario überhaupt aufgetreten; betroffen war in diesem Fall auch nur ein einzelnes Gerät.

Durch die Verdopplung der Simulationszeit von 900s auf 1800s verdoppelt sich ebenfalls auch nahezu die mittlere Zahl der Geräte, die im Verlauf der Simulation mindestens einmal auf

die maximale Sendeleistung von 100mW wechseln. Dies hat jedoch keinen nennenswerten Einfluss auf die Häufigkeiten der gewählten Sendeleistungen zum Ende der Simulationen.

8.4.4 Dynamische Szenarien mit Smallscale-Fading

Dieser Abschnitt beschreibt Ergebnisse der Simulationsläufe, in denen beide Arten von Dynamik modelliert wurden, sowohl Smallscale-Fading als auch Gerätemobilität.

Die Betrachtung von Eigenschaften wie der Nachbarzahl der Geräte oder den resultierenden Sendeleistungszuweisungen liefert hier keine neuen Erkenntnisse, vielmehr sind hier die gleichen Auswirkungen der Dynamik auf die unterschiedlichen Varianten der Sendeleistungsregelung zu beobachten, die in den vergangenen Abschnitten bereits detailliert beschrieben wurden. An dieser Stelle wird daher abschließend nur die Netzkapazität für verschiedene Varianten des Verfahrens und bei unterschiedlicher Dynamik betrachtet.

Wie aufgrund der im Vorangegangenen beschriebenen Beobachtungen zu erwarten ist, verhalten sich die Varianten des Verfahrens, die entweder auf das versuchsweise Reduzieren der Sendeleistungen oder auf die Redundanzsignalisierung verzichten, kaum anders als die Standardvariante.

Die in Abschnitt 8.4.2.2 für statische Szenarien bereits beschriebenen Unterschiede zwischen der Standardvariante und der Max-Degree-Variante des Verfahrens sind auch in dynamischen Szenarien zu beobachten: Für alle verwendeten Werte von $|\mu|$ (0,2m/s, 0,5m/s und 1m/s) ist mit der Max-Degree-Variante ein signifikant höherer Anteil der Ströme erfolglos. Wie ebenfalls in Abschnitt 8.4.2.2 beschrieben, hängt dies aber auch damit zusammen, dass die Max-Degree-Variante bei gleichem Zielintervall für die Nachbarzahlen weniger Links etabliert als die Standardvariante. Mit $k_{\min} = 10$ und $k_{\max} = 18$ verringert sich der Anteil erfolgloser Ströme für die Max-Degree-Variante, und während für $|\mu| = 0,2\text{m/s}$ auch mit diesen Parametern noch signifikant mehr Ströme erfolglos sind als mit der Standardvariante (wobei der Unterschied geringer ausfällt als in statischen Szenarien), gleicht sich dieser Wert für $|\mu| = 0,5\text{m/s}$ und $|\mu| = 1\text{m/s}$ auf das Niveau der anderen Varianten an.

Wird dieses Zielintervall mit den anderen Varianten des Verfahrens benutzt, verringert sich der Anteil erfolgloser Datenströme bei höherer Dynamik übrigens auch, allerdings nicht signifikant. Außerdem wurde bereits festgestellt, dass ab $|\mu| = 0,5\text{m/s}$ kaum noch ein Kapazitätsgewinn durch Erhöhung des Spatial-Reuse erzielbar ist. Diese Annäherung zwischen der Standardvariante und der Max-Degree-Variante des Algorithmus bei steigender Dynamik liegt also im geringen erzielbaren Kapazitätsgewinn begründet.

8.5 Fazit

Die Ergebnisse der im Vorangegangenen beschriebenen Untersuchungen zeigen zunächst, dass das Sammeln lokaler Topologieinformation und deren Nutzung durch das Routingpro-

protokoll zu einer erheblichen Steigerung der Leistungsfähigkeit eines drahtlosen Netzes führen können. Weiterhin kann das in Kapitel 7 beschriebene Verfahren diese Information zur dynamischen Regelung der Sendeleistungen der einzelnen Geräte nutzen, um einen möglichst hohen Spatial-Reuse zu erzielen. Vorkonfigurierte Sendeleistungen können in dieser Hinsicht hingegen nur günstig in Bezug auf eine konkrete Gerätedichte (im Verhältnis zum Sendebereich eines Geräts) sein.

Die Netzkapazität ist beim Einsatz des Algorithmus zur Sendeleistungsregelung höher oder zumindest vergleichbar zu dem Fall, dass alle Geräte mit maximaler Sendeleistung betrieben werden, wenn drei Bedingungen erfüllt sind:

1. Das Netz besteht aus einer großen Anzahl von Geräten. (Diese Bedingung wurde bereits in Kapitel 6 identifiziert.)
2. Die aus der Mobilität der Geräte resultierende Dynamik ist moderat.
3. Die räumliche Ausdehnung des Netzes ist groß im Vergleich zur Sendereichweite eines Geräts bei maximaler Sendeleistung.

Die letzte Bedingung trifft für das hier modellierte IEEE-802.11-Protokoll zu. Es ist angesichts der Ergebnisse aus Kapitel 6 jedoch nicht auszuschließen, dass mit verbesserten Protokollen für die Verbindungsschicht in drahtlosen Multihop-Netzen auch dann eine Kapazitätssteigerung durch Erhöhung des Spatial-Reuse erzielbar ist, wenn diese Bedingung nicht erfüllt ist.

Die Eigenheiten der Varianten des vorgeschlagenen Verfahrens (einschließlich der Max-Degree-Variante) machen sich in Abhängigkeit von den Szenarioeigenschaften sehr unterschiedlich bemerkbar. Die Kooperation zur Umsetzung des Nearest-Neighbours-Ansatzes als eine der wesentlichen Neuerungen gegenüber anderweitig bereits vorgeschlagenen Max-Degree-Verfahren hat sich insgesamt aber als wichtiges Mittel herausgestellt, um eine möglichst homogene Verteilung von Sendeleistungen zu erzielen. Der in Abschnitt 4.2.1.1 beschriebene Nachteil des Max-Degree-Ansatzes wird dadurch deutlich, dass einige Geräte mit maximaler Leistung senden, da sie nicht in der Lage sind, eine hinreichend hohe Zahl von Links aufzubauen, und dass weitere Geräte dafür unnötig hohe Sendeleistungen benötigen. In dynamischeren Szenarien erhöht sich die Deutlichkeit dieser Resultate noch im Vergleich zu statischen Szenarien.

Weiterhin wurden in Abschnitt 7.5 mehrere Mechanismen vorgestellt, die die Verringerung von Sendeleistungen regeln. Zunächst haben die Untersuchungen in diesem Kapitel gezeigt, dass mit der einfachen Variante, wie sie in [GdWMJ03] und [GdWMJ05] vorgestellt wurde, ein unakzeptabel hoher Anteil der Geräte zu viele Links hat. Die Signalisierung redundanter Links ist in statischen Szenarien (die auch in dem Sinn statisch sind, dass die Linkzustände keinen Fluktuationen unterworfen sind) ebenfalls nicht gut geeignet, um die Nachbarzahlen eines ausreichend hohen Anteils aller Geräte in den vorgegebenen Zielbereich zu bringen. Das lässt sich in diesem Fall durch das versuchsweise Reduzieren der Sendeleistungen erreichen,

welches die Redundanzsignalisierung überflüssig macht. Ist allerdings eine gewisse Dynamik vorhanden, entweder durch die Mobilität der Geräte oder durch die durch Smallscale-Fading verursachte Fluktuation von Linkzuständen, so ist die Redundanzsignalisierung bereits ausreichend. Daher kann in diesen Fällen wiederum auf den vergleichsweise komplizierten Mechanismus des versuchsweisen Reduzierens der Sendeleistungen verzichtet werden, zumal dieser auch zu einer höheren Zahl unnötiger Sendeleistungsanpassungen führt.

9 Zusammenfassung und Ausblick

9.1 Zusammenfassung

Die vorliegende Arbeit hat den Einfluss der Sendeleistungen der beteiligten Geräte auf die Kapazität drahtloser Netze untersucht, in denen ein gemeinsamer Übertragungskanal genutzt wird. In diesen Netzen ermöglicht eine Verringerung von Sendeleistungen eine Steigerung des *Spatial-Reuse*: Da das Signal eines Senders weniger andere Geräte mit einer nennenswerten Stärke erreicht, können mehr Sender gleichzeitig (erfolgreich) aktiv sein. Dadurch lässt sich unter gewissen Umständen eine Kapazitätssteigerung erzielen.

Zur Untersuchung dieses Phänomens auf abstrakterer Ebene ist in Kapitel 5 ein neues Verfahren zur Abschätzung der Transportkapazität von Topologien drahtloser Netze eingeführt worden. Es unterscheidet sich darin von anderen, bereits zuvor zum gleichen Zweck entwickelten Verfahren, dass es die Bewertung von Topologien mit einer hohen Anzahl von Stationen und Datenströmen in einer akzeptablen Rechenzeit ermöglicht.

Mit Hilfe dieses Verfahrens wurden in Kapitel 6 zunächst Topologien untersucht, die aus der Zuweisung einer einheitlichen Sendereichweite zu jeder Station entstehen. Daraus konnten Erkenntnisse hinsichtlich der Frage gewonnen werden, in welchen Szenarien eine Erhöhung des *Spatial-Reuse* überhaupt eine Kapazitätssteigerung ermöglicht. Dann wurden zwei weitere Topologieklassen untersucht. Dies waren zum einen die *Max-Degree-Topologien*, die sich aus der Idealisierung einer bereits bekannten Gruppe von Verfahren zur Sendeleistungsregelung ergibt, die die maximale Nachbarzahl aller Stationen im Netz beschränken. Die in gewisser Weise ähnlichen *Nearest-Neighbours-Topologien* basieren auf der Idee, dass jede Station nach Möglichkeit eine minimale Anzahl von Nachbarn in der Topologie hat. Es konnte gezeigt werden, dass die *Nearest-Neighbours-Topologien* den *Max-Degree-Topologien* vorzuziehen sind. Schließlich wurde gezeigt, dass auch andere Topologieklassen, die in der Forschung zur Sendeleistungsregelung in drahtlosen Netzen zu finden sind, keine Steigerung der Transportkapazität gegenüber *Nearest-Neighbours-Topologien* ermöglichen.

Aufbauend auf diesen Ergebnissen wurde in Kapitel 7 ein neues Verfahren vorgestellt, das den Geräten eines drahtlosen Netzes Sendeleistungen gemäß dem *Nearest-Neighbours-Prinzip* zuweist. Das Verfahren ist vollkommen dezentral, da es weder eine zentrale Instanz im laufenden Betrieb benötigt, noch eine Abstimmung der Gerätekonfigurationen erfordert. Der Signalisierungs-overhead ist sehr gering: Die Signalisierungsrahmen von Mechanismen zur Erkennung zuverlässiger Links, die in solchen Netzen ohnehin benötigt werden, wie im Rahmen der vorliegenden Arbeit ebenfalls gezeigt wurde, müssen lediglich um eine relativ geringe, konstante

Anzahl von Bits erweitert werden.

Das vorgestellte Verfahren ist außerdem fast bedingungslos praktisch einsetzbar, die einzige Voraussetzung ist die, dass zumindest ein Teil der Geräte die Einstellung der Sendeleistung überhaupt erlaubt. Damit unterscheidet sich das präsentierte Verfahren von vielen anderen bereits vorgestellten Verfahren, die auf geographische Daten zurückgreifen, die nur in statischen Szenarien einsetzbar sind oder die spezielle Dienste der Verbindungsschicht voraussetzen.

In Kapitel 8 wurde dieses verteilte Verfahren in Simulationen untersucht, in denen aktuelle Protokolle für die Funkübertragung und das Routing in drahtlosen Netzen modelliert wurden. Die modellierten Umgebungen waren unterschiedlicher Art, was die Mobilität der Geräte und die Berücksichtigung entfernungsunabhängiger Signalstärkeschwankungen angeht. Weiterhin wurden auch verschiedene Varianten des Verfahrens gesondert betrachtet. Dadurch ergibt sich ein genaues Bild dessen, welche Kombinationen von Mechanismen in bestimmten Umgebungen sinnvoll sind. Grundsätzlich zeigen die Ergebnisse aber auch, dass die Standardvariante des Verfahrens generell zu günstigen Netzeigenschaften führt, sofern drei Bedingungen erfüllt sind: Das Netz muss aus einer *großen Anzahl von Geräten* bestehen, die *nicht zu schnell* bewegt werden und die *über einen großen Bereich verteilt* sind (im Verhältnis zur Fläche, die ein Gerät mit maximaler Sendeleistung abdeckt). Damit ist die Eignung des Verfahrens für die in Kapitel 3 beschriebenen Anwendungsszenarien in der Regel gegeben, zumal die Möglichkeit besteht, dass die letzte Bedingung mit verbesserten Protokollen für die Verbindungsschicht in drahtlosen Netzen hinfällig wird.

Dass eine Verringerung der Sendeleistungen nur bei höherer Gerätezahl durch eine Steigerung des Spatial-Reuse zu einer Kapazitätssteigerung führt, wurde bereits in Kapitel 6 festgestellt. Eine kritische Masse von Geräten muss aber vorhanden sein, wenn das Netz überhaupt Bestand haben und für die Teilnehmer interessant sein soll. Bei hoher Mobilität der Geräte ist es weiterhin sinnvoller, dass die Geräte mit maximaler Leistung senden, um die Lebensdauer der Links zu maximieren und damit unnötig häufige Verbindungsunterbrechungen zu vermeiden. In Kapitel 3 wurde jedoch dargelegt, warum in den angedachten Anwendungsszenarien keine hohe Dynamik zu erwarten ist.

Ist die räumliche Ausdehnung des Netzes schließlich relativ klein im Verhältnis zu den maximal erzielbaren Sendereichweiten, so wird durch eine Verringerung der Sendeleistungen ausgehend vom Maximalwert zunächst nur erreicht, dass die Länge der Kommunikationspfade steigt, ohne dass der Spatial-Reuse dadurch ebenfalls erhöht würde. Während in Kapitel 6 festgestellt wurde, dass ein Betrieb der Geräte im Bereich sehr viel geringerer Sendeleistungen eine höhere Netzkapazität zulässt, ist dies in den modellierten Szenarien nicht der Fall, weil hier auch der durch eine erhöhte Anzahl von Einzelübertragungen entstehende Overhead auf Verbindungsebene ins Gewicht fällt. Daher ist nicht auszuschließen, dass mit Protokollen, in denen dieser Overhead geringer ist, eine Maximierung des Spatial-Reuse auch in diesen Umgebungen lohnenswert ist.

9.2 Ausblick

Im Laufe der Arbeit sind verschiedene Fragestellungen aufgeworfen worden, die außerhalb des behandelten Kernthemas liegen und deshalb nicht detailliert behandelt werden konnten, die jedoch interessante Ausgangspunkte für weiterführende Arbeiten sind.

In der vorliegenden Arbeit ist die Regelung der Sendeleistung als Mechanismus betrachtet worden, der Information über die lokale Umgebung eines Geräts, die die Linkerkennung liefert, als Eingabe weiterverarbeitet. Wie in Abschnitt 8.3.1.1 angesprochen wurde, sind Linkerkennung und Sendeleistungsregelung in gewisser Hinsicht aber komplementär, da sich die Netztopologie auch durch die dynamische Anpassung der Linkerkennungsparameter, die in der vorliegenden Arbeit bewusst unberücksichtigt geblieben ist, beeinflussen lässt (beispielsweise führen im Allgemeinen sowohl höhere Sendeleistungen als auch höhere Fehlertoleranzen bei der Linkerkennung zu mehr Links). Daher sollte eine engere Verzahnung dieser beiden Mechanismen angestrebt werden.

Die effiziente Unterstützung von Geräten, die unterschiedliche Signalkodierungen beherrschen, ist in drahtlosen Multihop-Netzen eine bislang wenig beachtete Herausforderung. Eine wichtige Voraussetzung zu deren Lösung ist die Kenntnis der Fehlerraten entlang von Links in Abhängigkeit der verwendeten Kodierung. Dieser Aspekt muss in der oben angesprochenen Verzahnung einer dazu fähigen Linkerkennung und dem in der vorliegenden Arbeit eingeführten verteilten Verfahren zur Sendeleistungsregelung ebenfalls berücksichtigt werden, da die Sendeleistungen nicht mehr nur das einfache Vorhandensein eines Links beeinflussen, sondern die maximale Datenrate, mit der eine Übertragung durchgeführt werden kann (so dass sie mit hoher Wahrscheinlichkeit erfolgreich ist).

Das vorgestellte Verfahren zur verteilten Sendeleistungsregelung ist ein Baustein für ein selbstkonfigurierendes drahtloses Netz, da die Sendeleistungen der Geräte transparent für deren Benutzer an die vorhandene Gerätedichte angepasst werden. Anstatt die konkrete Sendeleistung eines Geräts vorzugeben, wird ein Zielintervall für die Anzahl bidirektionaler Links eingestellt, was den Vorteil hat, dass eine günstige Wahl dieses Parameters unabhängig von der Gerätedichte ist. Wie bereits erwähnt, müssen allerdings drei Bedingungen erfüllt sein, damit sich daraus keine Nachteile gegenüber der Verwendung maximaler Sendeleistungen ergeben können. Weitere Arbeit kann deshalb auf die Erarbeitung eines Mechanismus abzielen, der Situationen erkennt, in denen eine dieser Bedingungen nicht erfüllt ist und den Geräten ggf. maximale Sendeleistungen unabhängig von deren Nachbarzahlen zuordnet. Des weiteren wäre es hilfreich, eine Verringerung der Netzkonnektivität zu erkennen und ihr in geeigneter Weise entgegenzuwirken.

Dabei stellt sich die Frage, welche zusätzliche Information die Geräte für solche Verbesserungen benötigen und mit welchem Aufwand deren Beschaffung verbunden ist. Sie kann wiederum durch eigens dafür vorgesehene Mechanismen beschafft werden, jedoch könnte auch eine geschickte Auswertung der Information weiterhelfen, die andere Protokolle ohnehin zur Verfügung stellen. Ein einfaches Beispiel für eine solche Interaktion wird in [RRH00] beschrieben: Anhand der Information eines proaktiven Link-State-Routingprotokolls erkennt

einer der dort vorgestellten Mechanismen, welche Links für die Konnektivität des Netzes wichtig sind und versucht, diese aufrechtzuerhalten. Eine so detaillierte Sicht der Netztopologie ist zweifellos ein mächtiges Werkzeug, aber der Overhead, der mit der Verteilung dieser Information an alle Netzteilnehmer verbunden ist, ist entsprechend hoch, weshalb proaktive Routingprotokolle als ungeeignet für große Netze gelten. Deshalb wäre zu erforschen, ob und wie weniger detaillierte Information dazu geeignet ist, die Entscheidungen eines Algorithmus zur Sendeleistungsregelung zu verbessern.

Literaturverzeichnis

- [Ari84] ARIKAN, ERDAL: *Some Complexity Results about Packet Radio Networks*. IEEE Transactions on Information Theory, IT-30(4):681–685, Juli 1984.
- [Bam98] BAMBOS, NICHOLAS: *Toward Power-sensitive Network Architectures in Wireless Communications: Concepts, Issues, and Design Aspects*. IEEE Personal Communications, 5:50–59, Juni 1998.
- [Bar01] BARTOSCH, LORENZ: *Generation of Colored Noise*. International Journal of Modern Physics C, 12(6):851–855, 2001.
- [BCF04] BORGONOVO, FLAMINIO, MATTEO CESANA und LUIGI FRATTA: *Broadcast Services and Topology Control in Ad Hoc Networks*. In: *Proc. IFIP TC6 / WG6.8 Conference on Mobile and Wireless Communication Networks (MW-CN'04)*, Seiten 407–418, November 2004.
- [BDEK02] BOSE, PROSENJIT, LUC DEVROYE, WILLIAM S. EVANS und DAVID G. KIRKPATRICK: *On the Spanning Ratio of Gabriel Graphs and beta-skeletons*. In: *Proc. 5th Latin American Symposium on Theoretical Informatics (LATIN)*, Seiten 479–493, April 2002.
- [Bet02] BETTSTETTER, CHRISTIAN: *On the Minimum Node Degree and Connectivity of a Wireless Multihop Network*. In: *Proc. 3rd ACM International Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc'02)*, Seiten 80–91, Juni 2002.
- [BFK⁺02] BAATZ, SIMON, MATTHIAS FRANK, CARMEN KUEHL, PETER MARTINI und CHRISTOPH SCHOLZ: *Bluetooth Scatternets: An Enhanced Adaptive Scheduling Scheme*. In: *Proc. IEEE INFOCOM 2002*, Seiten 782–790. IEEE, Juni 2002.
- [BJ02] BORBASH, STEVE A. und ESTHER H. JENNINGS: *Distributed Topology Control Algorithm for Wireless Multihop Networks*. In: *Proc. IEEE International Joint Conference on Neural Networks (IJCNN)*, Seiten 355–360, Mai 2002.

- [BLRS03] BLOUGH, DOUGLAS, MAURO LEONCINI, GIOVANNI RESTA und PAOLO SANTI: *The k-Neigh Protocol for Symmetric Topology Control in Ad Hoc Networks*. In: *Proc. 4th ACM International Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc'03)*, Seiten 141–152, Juni 2003.
- [Blu04] BLUETOOTH SIG: *Specification of the Bluetooth System – Version 2.0 + EDR*, November 2004.
- [Bré79] BRÉLAZ, DANIEL: *New Methods to Color Vertices of a Graph*. *Communications of the ACM (CACM)*, 22(4):251–256, April 1979.
- [BvRWZ04] BURKHART, MARTIN, PASCAL VON RICKENBACH, ROGER WATTENHOFER und AARON ZOLLINGER: *Does Topology Control Reduce Interference?* In: *Proc. 5th ACM International Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc'04)*, Seiten 9–19, Mai 2004.
- [BW02] BETTSTETTER, CHRISTIAN und CHRISTIAN WAGNER: *The Spatial Node Distribution of the Random Waypoint Mobility Model*. In: *Proc. 1st German Workshop on Mobile Ad-Hoc Networks (WMAN'02)*, Seiten 41–58, Ulm, Germany, März 2002.
- [CE04] CERPA, ALBERTO und DEBORAH ESTRIN: *ASCENT: Adaptive Self-Configuring sEnsor Networks Topologies*. *IEEE Transaction on Mobile Computing*, 3(3):272–285, Juli 2004.
- [Cis] Cisco Aironet 350 Series Client Adapters, data sheet. http://www.cisco.com/en/US/products/hw/wireless/ps4555/products_data_sheet09186a0080088828.html.
- [DRL03] D'SOUZA, RAISSA M., SHARAD RAMANATHAN und DUNCAN TEMPLE LANG: *Measuring Performance of Ad Hoc Networks Using Timescales for Information Flow*. In: *Proc. IEEE INFOCOM*, Seiten 1564–1574, April 2003.
- [DRWT97] DUBE, ROHIT, CYNTHIA D. RAIS, KUANG-YEH WANG und SATISH K. TRIPATHI: *Signal Stability based Adaptive Routing (SSA) for Ad-Hoc Mobile Networks*. *IEEE Personal Communication*, Seiten 36–45, Februar 1997.
- [DTH02] DOUSSE, OLIVIER, PATRICK THIRAN und MARTIN HASLER: *Connectivity in ad-hoc and hybrid networks*. In: *Proc. IEEE INFOCOM*, Seiten 1079–1088, Juni 2002.
- [EKCD00] ELBATT, TAMER A., SRIKANTH V. KRISHNAMURTHY, DENNIS CONNORS und SON DAO: *Power Management for Throughput Enhancement in Wireless Ad-Hoc Networks*. In: *Proc. IEEE International Conference on Communications (ICC)*, Seiten 1506–1513, Juni 2000.

- [EPY97] EPPSTEIN, DAVID, MICHAEL S. PATERSON und FRANCES F. YAO: *On Nearest Neighbor Graphs*. Discrete & Computational Geometry, 17(3):263–282, April 1997.
- [Eur98] EUROPEAN TELECOMMUNICATIONS STANDARDS INSTITUTE: *Digital cellular telecommunications system (Phase 2); Radio subsystem link control (GSM 05.08 version 4.22.1)*, März 1998. European Telecommunication Standard 300 578.
- [Eur99] EUROPEAN TELECOMMUNICATIONS STANDARDS INSTITUTE: *Digital cellular telecommunications system (Phase 2); Radio transmission and reception (GSM 05.05 version 4.23.1)*, Dezember 1999. European Telecommunication Standard 300 577.
- [FLZ⁺05] FU, ZHENGHUA, HAIYUN LUO, PETROS ZERFOS, SONGWU LU, LIXIA ZHANG und MARIO GERLA: *The Impact of Multihop Wireless Channel on TCP Performance*. IEEE Transactions on Mobile Computing, 4(2):209–221, März-April 2005.
- [FV06] FALL, KEVIN und KANNAN VARADHAN (Herausgeber): *The Ns Manual*. The VINT Project, UC Berkeley, LBL, USC/ISI, and Xerox PARC, März 2006.
- [GC04] GOMEZ, JAVIER und ANDREW T. CAMPBELL: *A Case for Variable-Range Transmission Power Control in Wireless Multihop Networks*. In: *Proc. IEEE INFOCOM*, März 2004.
- [GdWFM02] GERHARZ, MICHAEL, CHRISTIAN DE WAAL, MATTHIAS FRANK und PETER MARTINI: *Link Stability in Mobile Wireless Ad Hoc Networks*. In: *Proceedings of the 27th Annual IEEE Conference on Local Computer Networks (LCN'02)*, Seiten 30–39, Tampa, FL, November 2002.
- [GdWFM04] GERHARZ, MICHAEL, CHRISTIAN DE WAAL, MATTHIAS FRANK und PETER MARTINI: *Influence of Transmission Power Control on the Transport Capacity of Wireless Multihop Networks*. In: *Proc. 15th IEEE International Symposium on Personal, Indoor and Mobile Radio Communications (PIMRC'04)*, Seiten 1016–1021, September 2004.
- [GdWM04] GERHARZ, MICHAEL, CHRISTIAN DE WAAL und PETER MARTINI: *How To Discover Optimal Routes In Wireless Multihop Networks*. In: *Proc. IFIP TC6 / WG6.8 Conference on Mobile and Wireless Communication Networks (MW-CN'04)*, Seiten 323–334, November 2004.
- [GdWMJ03] GERHARZ, MICHAEL, CHRISTIAN DE WAAL, PETER MARTINI und PAUL JAMES: *A Cooperative Nearest Neighbours Topology Control Algorithm for*

- Wireless Ad Hoc Networks*. In: *Proc. 12th International Conference on Computer Communications and Networks (ICCCN'03)*, Seiten 412–417, Oktober 2003.
- [GdWMJ05] GERHARZ, MICHAEL, CHRISTIAN DE WAAL, PETER MARTINI und PAUL JAMES: *A Cooperative Nearest Neighbours Topology Control Algorithm for Wireless Ad Hoc Networks*. *Telecommunication Systems*, 28(3-4):317–331, März 2005.
- [GK78] GOWDA, K. CHIDANANDA und G. KRISHNA: *Agglomerative Clustering Using the Concept of Mutual Nearest Neighbourhood*. *Pattern Recognition*, 10(2):105–112, 1978.
- [GK98] GUPTA, PIYUSH und P. R. KUMAR: *Critical Power for Asymptotic Connectivity in Wireless Networks*. In *Stochastic Analysis, Control, Optimization and Applications: A Volume in Honor of W. H. Fleming*. W. M. McEneaney, G. Yin, Q. Zhang, eds. Birkhauser, 1998.
- [GK00] GUPTA, PIYUSH und P. R. KUMAR: *The Capacity of Wireless Networks*. *IEEE Transactions on Information Theory*, IT-46(2):388–404, März 2000.
- [Gra00] GRACE, KEVIN H.: *Mobile Mesh Link Discovery Protocol*. Internet-Draft, expired, IETF, September 2000.
- [GS69] GABRIEL, KUNO RUBEN und ROBERT R. SOKAL: *A new statistical approach to geographic variation analysis*. *Systematic Zoology*, 18:259–278, 1969.
- [GT01] GROSSGLAUSER, MATTHIAS und DAVID TSE: *Mobility Increases the Capacity of Ad-hoc Wireless Networks*. In: *Proc. IEEE INFOCOM*, Seiten 1360–1369, April 2001.
- [HM04a] HEKMAT, RAMIN und PIET VAN MIEGHEM: *Interference in Wireless Multi-hop Ad-hoc Networks and its Effect on Network Capacity*. *ACM Wireless Networks*, 10(4):389–399, Juli 2004.
- [HM04b] HEKMAT, RAMIN und PIET VAN MIEGHEM: *Study of Connectivity in Wireless Ad-hoc Networks with an Improved Radio Model*. In: *Proc. 2nd Workshop on Modeling and Optimization in Mobile, Ad Hoc and Wireless Networks*, März 2004.
- [HP05] HAENGGI, MARTIN und DANIELE PUCCINELLI: *Routing in Ad Hoc Networks: A Case for Long Hops*. *IEEE Communications Magazine*, 43(10):93–101, Oktober 2005.
- [HS88] HAJEK, BRUCE und GALEN SASAKI: *Link Scheduling in Polynomial Time*. *IEEE Transactions on Information Theory*, 34(5):910–917, September 1988.

- [Hu93] HU, LIMIN: *Topology Control for Multihop Packet Radio Networks*. IEEE Transactions on Communications, COM-41(10):1474–1481, Oktober 1993.
- [IEE99] IEEE LAN/MAN STANDARDS COMMITTEE: *Wireless LAN Medium Access Control (MAC) and Physical Layer (PHY) Specifications*, 1999. ANSI/IEEE Std. 802.11.
- [IEE02] IEEE LAN/MAN STANDARDS COMMITTEE: *Wireless Medium Access Control (MAC) and Physical Layer (PHY) Specifications for Wireless Personal Area Networks (WPANs)*, 2002. ANSI/IEEE Std. 802.15.1.
- [Int94] INTERNATIONAL ORGANIZATION FOR STANDARDIZATION / INTERNATIONAL ELECTROTECHNICAL COMMISSION: *Information technology – Open Systems Interconnection – Basic Reference Model: The Basic Model*, 1994. International Standard ISO/IEC 7498-1.
- [Jaf81] JAFFE, JEFFREY M.: *Bottleneck Flow Control*. IEEE Transactions on Communications, COM-29(7):954–962, Juli 1981.
- [JM96] JOHNSON, DAVID B. und DAVID A. MALTZ: *Dynamic Source Routing in Ad Hoc Wireless Networks*. In: IMIELINSKI und KORTH (Herausgeber): *Mobile Computing*, Band 353. Kluwer Academic Publishers, 1996.
- [JPPQ03] JAIN, KAMAL, JITENDRA PADHYE, VENKATA N. PADMANABHAN und LILI QIU: *Impact Of Interference On Multi-hop Wireless Network Performance*. In: *Proc. Ninth Annual International Conference on Mobile Computing and Networking (ACM MobiCom’03)*, Seiten 66–80, September 2003.
- [JS03] JUN, JANGEUN und MIHAIL L. SICHITIU: *The Nominal Capacity of Wireless Mesh Networks*. IEEE Wireless Communications, 10:8–14, Oktober 2003.
- [Kah77] KAHN, ROBERT E.: *The Organization of Computer Resources into a Packet Radio Network*. IEEE Transactions on Communications, COM-25(1):169–178, Januar 1977.
- [KK03] KAWADIA, VIKAS und P. R. KUMAR: *Clustering and Power Control in Ad Hoc Networks*. In: *Proc. IEEE INFOCOM*, April 2003.
- [KKW03] KUBISCH, MARTIN, HOLGER KARL und ADAM WOLISZ: *Distributed Algorithms for Transmission Power Control in Wireless Sensor Networks*. In: *Proc. IEEE Wireless Communications and Networking Conference (WCNC’03)*, März 2003.
- [KMT98] KELLY, FRANK P., AMAN K. MAULLOO und DAVID K. H. TAN: *Rate control in communication networks: shadow prices, proportional fairness and stability*. Journal of the Operational Research Society, 49:237–252, 1998.

- [KN03] KODIALAM, MURALI und THYAGA NANDAGOPAL: *Characterizing Achievable Rates in Multi-hop Wireless Networks: The Joint Routing and Scheduling Problem*. In: *Proc. Ninth Annual International Conference on Mobile Computing and Networking (ACM MobiCom'03)*, Seiten 42–54, September 2003.
- [KS78] KLEINROCK, LEONARD und JOHN SILVESTER: *Optimum Transmission Radii For Packet Radio Networks or Why Six Is a Magic Number*. In: *Conference Record, National Telecommunications Conference*, Seiten 4.3.1–4.3.5, Dezember 1978.
- [KS87] KLEINROCK, LEONARD und JOHN SILVESTER: *Spatial Reuse in Multihop Packet Radio Networks*. *Proc. of the IEEE*, 75(1):156–166, Januar 1987.
- [LH99] LIANG, B. und Z. HAAS: *Predictive distance-based mobility management for PCS networks*. In: *Proc. IEEE INFOCOM*, Seiten 1377–1384, New York, März 1999. IEEE.
- [LHS03] LI, NING, JENNIFER C. HOU und LUI SHA: *Design and Analysis of an MST-Based Topology Control Algorithm*. In: *Proc. IEEE INFOCOM*, April 2003.
- [LL02] LIU, JILEI und BAOCHUN LI: *MobileGrid: Capacity-aware Topology Control in Mobile Ad Hoc Networks*. In: *Proc. 11th International Conference on Computer Communications and Networks (ICCCN'02)*, Seiten 570–574, Oktober 2002.
- [LNG03] LI, ZHIFEI, SUKUMAR NANDI und ANIL K. GUPTA: *Study of IEEE 802.11 Fairness and its Interaction with Routing Mechanism*. In: *Proc. IFIP International Conference on Mobile and Wireless Communication Networks (MWCN'03)*, Oktober 2003.
- [LNG04] LI, ZHIFEI, SUKUMAR NANDI und ANIL K. GUPTA: *Improving MAC Performance in Wireless Ad Hoc Networks Using Enhanced Carrier Sensing (ECS)*. In: *NETWORKING, Networking Technologies, Services, and Protocols; Performance of Computer and Communication Networks; Mobile and Wireless Communication, Third International IFIP-TC6 Networking Conference, Athens, Greece, May 9-14, 2004, Proceedings*, Seiten 600–612, 2004.
- [LW01a] LI, XIANG-YANG und PENG-JUN WANG: *Constructing Minimum Energy Mobile Wireless Networks*. In: *Proc 2nd ACM International Symposium on Mobile Ad Hoc Networking and Computing*, Seiten 283–286, Oktober 2001.
- [LW01b] LI, XIANG-YANG und PENG-JUN WANG: *Constructing Minimum Energy Mobile Wireless Networks*. *ACM SIGMOBILE Mobile Computing and Communications Review*, 5(4):55–67, Oktober 2001.

- [LWS04] LI, XIANG-YANG, YU WANG und WEN-ZAHN SONG: *Applications of k-Local MST for Topology Control and Broadcasting in Wireless Ad Hoc Networks*. IEEE Transactions on Parallel and Distributed Systems, 15(12):1057–1069, Dezember 2004.
- [MBH01] MONKS, JEFFREY P., VADUVUR BHARGAVAN und WEN-MEI W. HWU: *A Power Controlled Multiple Access Protocol for Wireless Packet Networks*. In: *Proc. IEEE INFOCOM*, Seiten 219–228. IEEE, 2001.
- [NHK05] NAHM, KITAE, AHMED HELMY und C.-C. JAY KUO: *TCP over Multihop 802.11 Networks: Issues and Performance Enhancement*. In: *Proc. 6th ACM International Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc'05)*, Seiten 277–287, Mai 2005.
- [NK85] NELSON, RANDOLPH und LEONARD KLEINROCK: *Spatial TDMA: A Collision-Free Multihop Channel Access Protocol*. IEEE Transactions on Communications, COM-33(9):934–944, September 1985.
- [NKSK02] NARAYANASWAMY, SWETHA, VIKAS KAWADIA, R. S. SREENIVAS und P. R. KUMAR: *Power Control in Ad-Hoc Networks: Theory, Architecture, Algorithm and Implementation of the COMPOW Protocol*. In: *Proc. European Wireless 2002 (EWC'02)*, Seiten 156–162, Februar 2002.
- [PB94] PERKINS, CHARLES E. und PRAVIN BHAGWAT: *Highly Dynamic Destination-Sequenced Distance-Vector Routing (DSDV) for Mobile Computers*. In: *Proc. ACM SIGCOMM*, Seiten 234–244, August 1994.
- [PB03] PEREVALOV, EUGENE und RICK BLUM: *Delay Limited Capacity of Ad hoc Networks: Asymptotically Optimal Transmission and Relaying Strategy*. In: *Proc. IEEE INFOCOM*, Seiten 1575–1582, April 2003.
- [PBRD03] PERKINS, CHARLES, ELIZABETH BELDING-ROYER und SAMIR DAS: *RFC 3561: Ad hoc On-Demand Distance Vector (AODV) Routing*, Juli 2003.
- [Pes04] PESCHLOW, PATRICK: *Entwurf und Bewertung von Verfahren zur Kapazitätsbestimmung von Topologien drahtloser Multihop-Netze*. Diplomarbeit, Rheinische Friedrich-Wilhelms-Universität Bonn, März 2004.
- [PGdWM04] PESCHLOW, PATRICK, MICHAEL GERHARZ, CHRISTIAN DE WAAL und PETER MARTINI: *An Efficient Transport Capacity Estimation Method for Wireless Multihop Network Topologies*. In: *Proc. 34th Annual Conference of the Gesellschaft für Informatik, Vol. 1 (2nd German Workshop on Mobile Ad-hoc Networks, WMAN'04)*, Seiten 90–94, September 2004.
- [Plu82] PLUMMER, DAVID C.: *RFC 826: An Ethernet Address Resolution Protocol*, November 1982.

- [PNS00] PUNNOOSE, RATISH J., PAVEL V. NIKITIN und DANIEL D. STANCIL: *Efficient Simulation of Ricean Fading within a Packet Simulator*. In: *IEEE Vehicular Technology Conference Fall*, Seiten 764–767, September 2000.
- [PS02a] PARK, SEUNG-JONG und RAGHUPATHY SIVAKUMAR: *Load-Sensitive Transmission Power Control in Wireless Ad-hoc Networks*. In: *Proc. IEEE Global Communications Conference (GLOBECOM)*, Seiten 42–46, November 2002.
- [PS02b] PARK, SEUNG-JONG und RAGHUPATHY SIVAKUMAR: *Quantitative Analysis of Transmission Power Control in Wireless Ad-hoc Networks*. In: *Proc. 31st International Conference on Parallel Processing Workshops (ICPP Workshops)*, Seiten 56–63, August 2002.
- [Ram97] RAMANATHAN, RAM: *A Unified Framework and Algorithm for (T/F/C)DMA Channel Assignment in Wireless Networks*. In: *Proc. IEEE INFOCOM*, Seiten 900–907, Kobe, Japan, April 1997.
- [Rap96] RAPPAPORT, THEODORE S.: *Wireless Communications – Principles & Practice*. Prentice Hall Communications Engineering and Emerging Technologies Series. Prentice Hall Professional Technical Reference, 1996.
- [RM98] RODOPLU, VOLKAN und TERESA H. MENG: *Minimum Energy Mobile Wireless Networks*. In: *Proc. IEEE International Conference on Communications (ICC)*, Seiten 1633–1639, Juni 1998.
- [RM99] RODOPLU, VOLKAN und TERESA H. MENG: *Minimum Energy Mobile Wireless Networks*. *IEEE Journal on Selected Areas in Communications*, 17(8):1333–1344, August 1999.
- [RRH00] RAMANATHAN, RAM und REGINA ROSALES-HAIN: *Topology Control of Multihop Wireless Networks using Transmit Power Adjustment*. In: *Proc. IEEE INFOCOM*, Seiten 404–413, 2000.
- [Rut91] RUTGERS, CHARLES L. HEDRICK: *An Introduction to IGRP*, August 1991. White paper.
- [SBV01] SANTI, PAOLO, DOUGLAS M. BLOUGH und FEODOR VAINSTEIN: *A Probabilistic Analysis for the Range Assignment Problem in Ad Hoc Networks*. In: *Proc. 2nd ACM International Symposium on Mobile Ad Hoc Networking and Computing (MobiHoc'01)*, Seiten 212–220, Oktober 2001.
- [SC99] SEN, ARUNABHA und JEFFREY M. CAPONE: *Scheduling in Packet Radio Networks - A New Approach*. In: *Proc. IEEE Globecom 1999*, Seiten 650–654, 1999.
- [Sha49] SHANNON, CLAUDE ELWOOD: *The Mathematical Theory of Information*. University of Illinois Press, 1949.

- [She95] SHENKER, SCOTT: *Fundamental Design Issues for the Future Internet*. IEEE Journal on Selected Areas in Communication, 13(7):1141–1149, September 1995.
- [SHS04] SUNDARESAN, KARTHIKEYAN, HUNG-YUN HSIEH und RAGHUPATHY SIVAKUMAR: *IEEE 802.11 over Multi-hop Wireless Networks: Problems and New Perspectives*. Ad Hoc Networks, 2(2):109–132, April 2004.
- [SM02] SUCEC, JOHN und I. MARSIC: *Clustering Overhead for Hierarchical Routing in Mobile Ad Hoc Networks*. In: *Proc. IEEE INFOCOM*, Seiten 1698–1706, Juni 2002.
- [SNK05] STOJMENOVIC, IVAN, AMIYA NAYAK und JOHNSON KURUVILA: *Design Guidelines for Routing Protocols in Ad Hoc and Sensor Networks with a Realistic Physical Layer*. IEEE Communications Magazine, 43(3):101–106, März 2005.
- [Tan03] TANENBAUM, ANDREW S.: *Computer Networks*. Prentice Hall Professional Technical Reference, 2003.
- [TK84] TAKAGI, HIDEAKI und LEONARD KLEINROCK: *Optimal Transmission Ranges for Randomly Distributed Packet Radio Terminals*. IEEE Transactions on Communications, COM-32(3):246–257, März 1984.
- [Tou80] TOUSSAINT, GODFRIED: *The Relative Neighborhood Graph of a Finite Planar Set*. Pattern Recognition, 12(4), 1980.
- [TSRK04] TRANTER, WILLIAM H., K. SAM SHANMUGAN, THEODORE S. RAPPA-PORT und KURT L. KOSBAR: *Principles of Communication Systems Simulation with Wireless Applications*. Prentice Hall Communications Engineering and Emerging Technologies Series. Prentice Hall Professional Technical Reference, 2004.
- [WC03] WOO, ALEC und DAVID CULLER: *Evaluation of Efficient Link Reliability Estimators for Low-Power Wireless Networks*. Technischer Bericht UCB//CSD-03-1270, U. C. Berkeley Computer Science Division, September 2003.
- [WLBW01] WATTENHOFER, ROGER, LI LI, PARAMVIR BAHL und YI-MIN WANG: *Distributed Topology Control for Power Efficient Operation in Multihop Wireless Networks*. In: *Proc. IEEE INFOCOM*, Seiten 1388–1397, 2001.
- [WTC03] WOO, ALEC, TERENCE TONG und DAVID CULLER: *Taming the Unerlying Challenges of Reliable Multihop Routing in Sensor Networks*. In: *Proc. ACM SenSys'03*, November 2003.

- [WYG⁺01] WU, XINRAN, CLEMENT YUEN, YAN GAO, HONG WU und BAOCHUN LI: *Fair Scheduling with Bottleneck Consideration in Wireless Ad-hoc Networks*. In: *Proc. of the 10th IEEE International Conference on Computer Communications and Networks (ICCCN'01)*, Seiten 568–572, Oktober 2001.
- [WZ04] WATTENHOFER, ROGER und AARON ZOLLINGER: *XTC: A Practical Topology Control Algorithm for Ad-Hoc Networks*. In: *Proc. 4th International IEEE Workshop on Algorithms for Wireless, Mobile, Ad Hoc and Sensor Networks (WMAN)*, April 2004.
- [XK04] XUE, FENG und P. R. KUMAR: *The Number of Neighbors Needed for Connectivity of Wireless Networks*. *Wireless Networks*, 10(2):169–181, März 2004.
- [XS01] XU, SHUGONG und TAREK SAADAWI: *Does the IEEE 802.11 MAC Protocol Work Well in Multihop Wireless Ad Hoc Networks?* *IEEE Communications Magazine*, Seiten 130–137, Juni 2001.
- [XS02] XU, SHUGONG und TAREK SAADAWI: *Performance Evaluation of TCP Algorithms in Multi-hop Wireless Packet Networks*. *Wireless Communications and Mobile Computing*, 2(1):85–100, Februar 2002.
- [YS03] YUEN, WING HO und CHI WAN SUNG: *On Energy Efficiency and Network Connectivity of Mobile Ad Hoc Networks*. In: *Proc. of the 23rd International Conference on Distributed Computing Systems (IEEE ICDCS'03)*, Seiten 38–45, Mai 2003.
- [YTCS99] YI, SZE-NAO, YU-CHEE TSENG, YUH-SHYAN CHEN und JANG-PING SHEU: *The Broadcast Storm Problem in a Mobile Ad Hoc Network*. In: *ACM Fifth Annual International Conference on Mobile Computing and Networking (MOBICOM'99)*, Seiten 151–162, Seattle, Washington, USA, August 1999.
- [ZK03] ZUNIGA, MARCO und BHASKAR KRISHNAMACHARI: *Optimal Transmission Radius for Flooding in Large Scale Sensor Networks*. In: *Proc. 23rd International Conference on Distributed Computing Systems Workshops (ICDCSW'03)*, Seiten 697–704, Mai 2003.