

Netzwerk-Messtechnik und Monitoring

Moderne Verkabelungs-Qualifizierer

**Analyse von Qualitätsschwankungen
bei Netzwerkdiensten**

**Mit Marktübersicht
Monitoring-Tools**



**Vernetzte Lösungen
für mobiles Computing**
Herausforderungen
bei der Datenerfassung

**Drahtlose
Präsentationen:**
BenQ InstaShow
unter der Lupe

**Künstli
im IT
M**

**Sonderdruck tde
Ausfallsicher und
hochverfügbar**

Zukunftsfähige und effiziente RZ-Verkabelung

Ausfallsicher und hochverfügbar

Heute schon an morgen denken – dieses Motto gilt für die RZ-Verkabelung mehr denn je. Schließlich verlangen die zunehmende Virtualisierung und Entwicklungen wie Software-Defined Networking (SDN), Cloud Computing oder Hyperscale-Datacenter von Verkabelungslösungen größere Bandbreiten und die Anpassung an künftige Anforderungen. Moderne Datacenter können sich keine Engpässe mehr erlauben und müssen Latenzen minimieren. Eine weitsichtige Planung mit ausreichenden Reserven ist technisch und betriebswirtschaftlich zwingend nötig.

Rechenzentren sind das Rückgrat der digitalisierten Informationsgesellschaft. Ihre Planung, den Bau und Betrieb regelt das europäische Regelwerk EN 50600. Auch die infrastrukturelle Ebene ist klar definiert und normiert: Betreiber und Planer erhalten mit den Normen EN 50600-2-4 zur Infrastruktur der Telekommunikationsverkabelung, der EN 50173-5 zu anwendungsneutralen Kommunikationskabelanlagen in Rechenzentren sowie der Norm EN 50174-2 zur Installation von Kommunika-

tionsverkabelungen in Gebäuden fundierte Hilfestellungen.

Dabei erlaubt die EN 50600-2-4 zwei Arten von Verkabelung im Rechenzentrum: Unter bestimmten Umständen möglich ist die Punkt-zu-Punkt-Verkabelung, allerdings nur, wenn sie innerhalb eines Racks oder zwischen maximal zwei nebeneinanderstehenden Racks erfolgt. Die zweite Variante ist die strukturierte Verkabelung gemäß EN 50173. Sie basiert auf der konventionellen Router-Architektur mit den

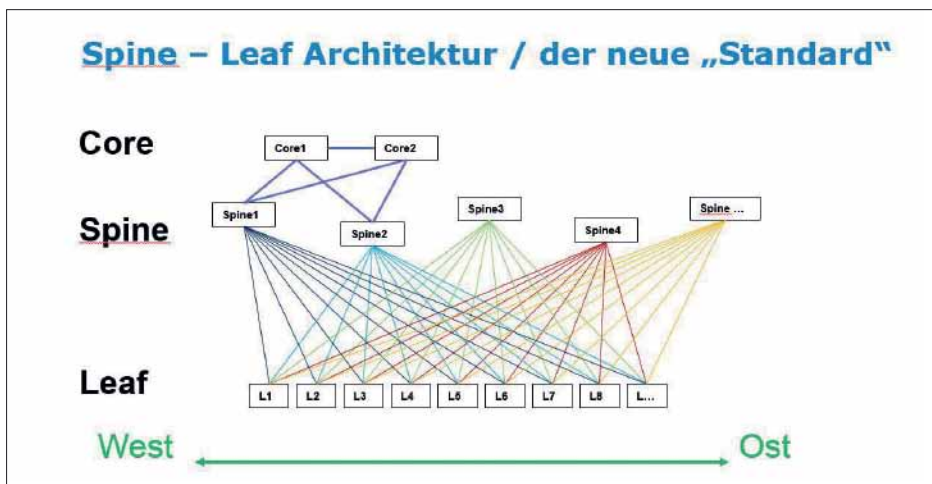
drei Ebenen Core-, Aggregation- und Access-Layer. Die Datenkommunikation erfolgt über externe Verbindungen in und aus dem Rechenzentrum.

Bei diesem North-South-Traffic steigen vor allem an den Schnittstellen zum externen Netz und den Hauptverteilern (North) die Anforderungen an die Bandbreite für die aktive und passive Technik. Ein direkter Austausch auf der East-West-Ebene findet hingegen nicht statt – ein Problem, dass im Fall von Firmenrechenzentren schnell zu Engpässen auf der Aggregationsebene führt. Dabei spielt die Veränderung der Datenflüsse eine entscheidende Rolle. Denn während bisher der Haupt-Traffic der externen zu den internen Datenverbindungen im Verhältnis von 4:1 stattfand, ist es heute umgekehrt. Der Großteil der Datenströme im Firmenrechenzentrum erfolgt auf einer Ebene (Aggregation/Layer 2). Ein wesentlicher Grund für die Zunahme des East-West-Traffics ist die Virtualisierung, bei der sich virtuelle Maschinen im Aggregations-Layer platzieren lassen und dort die gesamte Bandbreite ausnutzen.

Mit Spine-Leaf den East-West-Traffic bändigen

An diese Stelle setzt das Konzept der Spine-Leaf-Architektur an: Entwickelt, um die East-West-Kommunikation zu unterstützen, verbinden vollvermaschte Kreuzverbindungen die Spine-Ebene (Aggregations-Layer) mit der Leaf-Ebene (Access-Layer) und verkürzen die Latenzzeiten erheblich. Zugleich bedingt die Architektur, dass bei einer Erweiterung des Netzwerks um neue Spine- oder Leaf-Switches jeder einzelne Leaf- oder Spine-Switch miteinander verbunden sein muss. Dabei entspricht die Anzahl der möglichen Verbindungen der Anzahl der implementierten Spine-Switches. Eine Begrenzung der Kapazität erfolgt durch die Anzahl der Spine-Ports.

Im Gegensatz zum Spanning-Tree-Protokoll, das bei hierarchischen Strukturen zum Einsatz kommt, nutzt die Spine-Leaf-Architektur für das Multipath Routing auf Layer 2 die Netzwerkprotokolle TRILL (TRansparent Interconnection of Lots of



Vollvermaschte Kreuzverbindungen verbinden die Spine-Ebene (Aggregations-Layer) mit der Leaf-Ebene (Access-Layer) und verkürzen die Latenzen erheblich.

Bild: TDE – Trans Data Elektronik



Das kompakte tML-LWL-Spine-Leaf-Mesh-Modul vereinfacht die Verkabelung in komplexen Spine-Leaf-Architekturen, indem es zwischen vier und 16 Kanäle/Faserpaare zusammenfasst und wieder aufteilt. Dies reduziert die Zahl der Anschlüsse um ein Viertel, Achtel oder Sechzehntel.

Bild: TDE – Trans Data Elektronik



Plug-and-Play-fähige Verkabelungssysteme integrieren inzwischen einen 32-Faser-MPO und können Unternehmen zukunftsfähige und investitionssichere Migrationsoptionen aktuell bis zu zweimal 400 GBit/s bieten.

Bild: TDE – Trans Data Elektronik

Links) und das im IEEE-Standard 802.1aq spezifizierte SPB (Shortest Path Bridging). Beide erkennen die optimale Verbindung und können so vorhandene Ressourcen gleichmäßig und umfassend nutzen. Es entstehen auch keine Engpässe wie bei der klassischen dreischichtigen Router-Topologie. Zudem punktet die vermaschte Netzwerk-Topologie mit hoher Ausfallsicherheit. Selbst der Ausfall eines kompletten Spine-Switches hat nur geringe Auswirkungen.

Doch auch wenn die Spine-Leaf-Architektur optimale Verbindungen schafft, birgt ihre Systematik entscheidende Herausforderungen: Zum einen erhöht sich im Gegensatz zur rein hierarchischen Struktur die Anzahl der benötigten Switches auf der Spine-Ebene. Zum anderen steigt die Zahl der notwendigen Verbindungen zwischen der Leaf- und der Spine-Ebene extrem – und damit das Verbindungschaos. Der Grund: Da sich die Mindestanzahl an Verbindungen aus der Anzahl der Spine- mal der Anzahl der Leaf-Switches ergibt, können sich schon schlanke Kreuzverbindungs-Topologien schnell zum komplexen und unübersichtlichen Durcheinander ausweiten.

Spine-Leaf-Mesh-Module als Lösung

An dieser Stelle kommt die paralleloptische Übertragung ins Spiel. Sie trägt dazu bei, die Anzahl der physischen Verbindungen zu reduzieren, indem sie die Kanäle/Faserpaare parallel schaltet. Dies gelingt durch die Breakout-Funktion der QSFP-

Transceiver, die die Zusammenfassung von Ports auf der aktiven Seite ermöglichen. Die Kanalaufteilung und -zuordnung erfolgt über Mesh-Module. Lösungen wie das tML-LWL-Spine-Leaf-Mesh-Modul verringern die Komplexität der Kreuzverbindungen, indem es zwischen vier und 16 Kanäle/Faserpaare zusammenfasst und wieder aufteilt.

Entsprechend lässt sich auch die Zahl der Anschlüsse um ein Viertel, Achtel oder Sechzehntel reduzieren. Ein Mesh-Modul kann bei vier Kanälen 16 Verbindungen abdecken. Dabei ist es im Gegensatz zu vergleichbaren Breakout-Modulen deutlich platzsparender. Netzwerkadministratoren sparen sich mit einem Modul acht Breakout-Module und 32 LC-Duplex-Patch-Kabel ein. Da sich diese Kanalsammenfassung und -aufteilung über Mesh-Module potenzieren lässt, können Netzwerkadministratoren das Verbindungschaos in vorgefertigte Strukturen lenken und direkt auf das Mesh-Modul patchen.

Dies vereinfacht die Verkabelung in komplexen Spine-Leaf-Architekturen. Und auch auf anderen Ebenen spielen Mesh-Module ihre Vorteile aus: Indem die aufwendige Verkabelung mit LC-Steckverbindern entfällt, sparen Unternehmen fast 90 Prozent Platz im Patch-Schrank gegenüber bisherigen Lösungen. Zugleich reduziert sich das Kabelvolumen in den Kabelwegen und im Patch-Bereich drastisch, was eine geringere Brandlast und mehr Energieeffizienz bedeutet.

Da Netzwerktechniker nur noch acht MPO-Patch-Cords verkabeln müssen, reduziert sich die Installationszeit. Außerdem ist die Dokumentation weniger komplex und dadurch einfacher nachzuvollziehen. Und nicht zuletzt können Unternehmen mit einer professionellen Lösung Kosteneinsparungen von bis zu 50 Prozent erzielen.

Entwicklung der Übertragungsgeschwindigkeiten

Ein weiterer Aspekt der zukunftsfähigen Verkabelung betrifft die Übertragungsbreiten: Während von 1983 bis 2010 neue Normen in relativ kurzen Abständen mit jeweils einer Verzehnfachung der Leistung erschienen, war es 2016 mit 25 GBit/s pro Kanal nur das 2,5-fache. Seit 2018 lassen sich bis zu 50 GBit/s je Kanal übertragen, wobei dies keine Beschleunigung der Übertragungsleistung darstellt, sondern durch die Modifikation des Signaltyps entsteht. Die sogenannte PAM4-Modulation ermöglicht die Verdoppelung der Übertragungsleistung. Sie hat Eingang in die Normierung IEEE 802.3bs (Dezember 2017) gefunden und unterstützt Übertragungsleistungen von bis zu 400 GBit/s. Der im Dezember 2018 verabschiedete Standard IEEE 802.3cd legt die Übertragung von 50/100/200 GBit/s fest. Auch an dieser Stelle lässt sich durch Duplizierung oder Vervielfachung der Übertragungsleistung pro Kanal – derzeit 50 Gigabit pro Kanal – eine höhere Übertragungsgeschwindigkeit erreichen. Nach

dieser Systematik agieren derzeit alle Normen, um die dringend gewünschten hohen Übertragungsleistungen erreichen zu können.

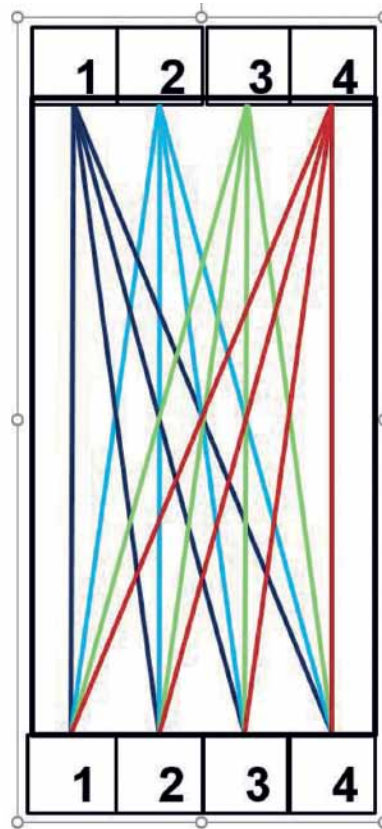
Wege zu den Terabit Mountains

Derzeit suchen die Normungsgremien Lösungen, um Übertragungsleistungen von 100 GBit/s über einen Kanal zu definieren. Durch die Vervielfältigung der Kanäle sollen sich 800 GBit beziehungsweise 1,6 TBit pro Sekunde über einen Port übertragen lassen. Doch auch abseits der Normierungsgremien ergeben sich interessante Entwicklungen: So haben Transceiver-Hersteller unter Rückgriff auf das elektronische Interface der IEEE einen SR8-Transceiver entwickelt, der Übertragungen von 50 GBit/s und Kanal und damit eine Übertragungsleistung von 400 GBit/s über Multimode als Übertragungsleistung ermöglicht.

Diese Option greift aktuell der Normentwurf IEE 802.3 cm auf, der im November 2019 in Kraft treten soll. Daneben kristallisiert sich im Rahmen dieser Norm auch die bidirektionale Übertragung als WDM-Übertragung heraus. Dabei lassen sich mit zwei Wellenlängen über bis zu vier Kanäle bis 400 GBit/s über acht Multimode-Fasern übertragen. Diese WDM-Anwendung bedient sich zur Leistungserhöhung der paralleloptischen Übertragung und verbindet erstmals beide Techniken miteinander. Bereits Standard ist die SWDM-Übertragung als reine WDM-Übertragung über Multimode mit 100 GBit/s über zwei Multimode-Fasern.

Die Qual der Übertragungsoptionen

Doch welche Übertragungsmedien sollten Netzwerktechniker nehmen, etwa beim Aufbau einer 100-Gigabit-Strecke? Klar ist: Das Angebot ist vielseitig. Für die strukturierte Verkabelung in Rechenzentren eignen sich QSFP-Transceiver mit Glasfaser über Multimode- oder Single-mode-Fasern und einer Übertragungsbreite von zwei bis acht Fasern. Da die Übertragung über einen Kanal begrenzt ist, erfordern hohe Bandbreiten zwingend die Bündelung der Kanäle. Diese lässt sich entweder durch die Parallelschaltung der



Schematische Darstellung eines Mesh-Moduls:
Die Abbildung zeigt die Zusammenfassung und Aufteilung von vier Kanälen/Faserpaaren.

Bild: TDE – Trans Data Elektronik

physischen Kanäle (Paralleloptik) oder durch WDM-Übertragung realisieren, bei der die Bündelung der unterschiedlichen Wellenlängen auf einem physikalischen Kanal erfolgt.

Derzeit lassen sich bis zu 16 Kanäle parallel übertragen, wobei die Parallelschaltung 32 Fasern je Port und die WDM-Anwendung zwei Fasern je Port nutzt. Auf der Seite der Anschlusstechniken stehen als Schnittstellen mit dem SFP DD oder QSFP DD inzwischen zwei Double-Density-Transceiver zur Verfügung, die auf gleicher Fläche wie normale SFP/SFP+/QSFP die zweifache Leistung erzielen. In Verbindung mit einem QSFP-DD-Transceiver haben Netzwerkadministratoren hier die Wahl zwischen sieben möglichen Schnittstellen: Zum einen ist dies der MPO-12-Faser-Stecker entweder mit einer Übertragung über vier Kanäle oder als paralleloptische BiDi-Übertragung. Alternativ dazu stehen ein MPO-16-Faser-Stecker mit acht Kanälen in einer Wellenlänge sowie ein MPO 24 mit zweimal vier Kanälen für

Breakout-Anwendungen zur Verfügung. Ebenfalls denkbar ist der neuartige und kompakte CS-Steckverbinder, von denen sich zwei in einem SR2-Transceiver für die Zweifachübertragung nutzen lassen. Weitere Optionen der Übertragung ergeben sich mit den neuen MDC- oder SN-Steckverbindern mit jeweils acht Fasern, wobei auch hier BiDi-Übertragungen möglich sind. Als reine WDM-Anbindung eignet sich die Verbindung über klassische LC-Duplex-Steckverbinder. Interessant dabei: Nur ein Anschluss basiert auf der Zweifaserververkabelung, alle anderen setzen auf die Paralleloptik. Für welche Option sollten sich Netzwerktechniker letztendlich entscheiden? In diesem Kontext lohnt eine Gegenüberstellung und ein Vergleich beider Techniken.

Paralleloptik oder WDM-Übertragung

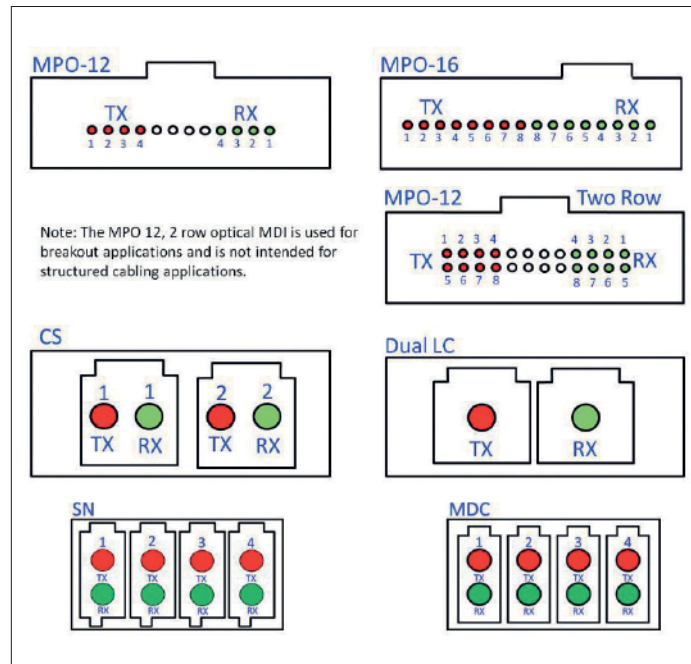
Im Gegensatz zur WDM-Übertragung, die mit nur zwei Fasern auskommt, birgt die paralleloptische Übertragung auf den ersten Blick den Nachteil des höheren Faserbedarfs. Dafür punktet sie an anderer Stelle: Bei der aktiven Technik ist sie günstiger als die WDM-Übertragung. Denn je nach Übertragungsgeschwindigkeit und Fasertyp können WDM-Transceiver bis zu fünfmal teurer sein als paralleloptische Transceiver mit identischer Leistung.

Ein weiterer Pluspunkt ist die Breakout-Funktion: Paralleloptische QSFP- oder CFP-Transceiver lassen sich so konfigurieren, dass sich über einen 40GbE-Port eines Switches vier Server mit 10 GBit/s ansteuern lassen. So lässt sich beispielsweise die 10- oder 25-GBit/s-Port-Dichte eines Switches von 48 Ports/HE auf 144 Ports/HE erhöhen. Dies führt zu einer Vervielfältigung der Port-Dichte auf der MPO-Seite bei einer gleichzeitig deutlichen Reduzierung des Energiebedarfs. Zugleich lässt sich die Bandbreite der recht kostspieligen und hochleistungsfähigen Transceiver optimal ausnutzen, und auch die Migration hin zu höheren Übertragungsraten wird unterstützt. Diese Optionen verlangt zwingend die physische Aufteilung der Kanäle – wahlweise über Fan-outs oder Breakout-Module. Bei WDM-Anwendungen sind diese Funktionen nicht möglich.

Beim Aufbau einer Netzwerkinfrastruktur auf Basis der parallel optischen Übertragung sollten Planer und Netzwerkadministratoren von Anfang an auf einen ausreichend hohen Faserbedarf und Reserven achten und künftige Erweiterungen berücksichtigen: Über sie lassen sich in Verbindung mit dem passenden Transceiver auch im Nachhinein noch WDM-Anwendungen realisieren. Sparen sie hingegen bei der Planung an den Fasern, lassen sich nachträgliche Veränderungen des Faserbedarfs nicht mehr umsetzen.

Hoher Faserbedarf im Rückraum

Plug-and-Play-fähige Verkabelungssysteme wie das modular aufgebaute tML-32-System integrieren inzwischen einen 32-Faser-MPO und können Unternehmen zukunfts-fähige und investitionssichere Migrationsoptionen bis zu aktuell zweimal 400 GBit/s bieten. Seit Dezember 2017 ist die 400GbE-Übertragung über Multimode-Fasern als derzeit einzige existierende Norm für den 32-Faser-MPO normiert. Aktuell gibt es dafür keine passenden Transceiver, weshalb die Realisierung derzeit über den 16-Faser-MPO erfolgt. Dieser ist



Schematische Zeichnung der möglichen Schnittstellen am Beispiel eines QSFP-DD-Transceivers: Sechs der sieben Anschlüsse basieren auf der Paralleloptik – nur einer auf der Zweifaserverkabelung (Dual LC).

Bild: TDE – Trans Data Elektronik

jedoch offiziell noch nicht normiert. Erst im November 2019 soll die Norm 802.3cm verabschiedet sein. An dieser Stelle zeigt sich die Diskrepanz zwischen der Normierung und den aktuellen Entwicklungen: Denn mit dem 16-Faser-MPO vom Hersteller TDE gibt es bereits qualitativ hochwertige Steckverbinder, und auch die passenden Transceiver stehen zur Verfügung. Damit lassen sich die 32 Fasern im Rückraum auf zwei 16-Faser-MPO umrouten. Da sich bei

modularen Systemplattformen alle Komponenten wiederverwenden lassen, ist auch ein Um-, Aus- und Rückbau sowie die Erweiterung jederzeit problemlos möglich.

Fazit

Unternehmen sollten bereits heute darauf achten, dass ihre Infrastruktur die neuen Anschlusstechniken wie den 16- oder 32-MPO-Faser-Stecker sowie die kompakten CS-, MDC- oder SN-Steckverbinder unterstützt. Nur so lässt sich eine tragfähige und langfristig zukunftssichere Verkabelung gewährleisten. Setzen Unternehmen zudem auf die parallel optische Übertragung, profitieren sie von geringeren Kosten bei der aktiven Technik, können komplexe Spine-

Leaf-Architekturen vereinfachen sowie Breakout-Funktionen und die WDM-Übertragung nutzen. Indem so alle Übertragungsszenarien möglich sind, trägt die Paralleloptik entscheidend dazu bei, eine investitionssichere und effiziente RZ-Verkabelung zu schaffen, und zwar ausfallsicher und hochverfügbar. Rainer Behr/jos

Rainer Behr ist Sales Consultant bei TDE Trans Data Elektronik, www.tde.de.

News

XDR soll mehr Sichtbarkeit und schnellere Analysen schaffen

Trend Micro: Detection and Response für E-Mail, Netzwerk, Endpunkte, Server und Cloud-Workloads

Der Sicherheitsanbieter Trend Micro hat vor Kurzem ein Konzept namens XDR in sein Produktportfolio eingeführt. Damit will man auf integrative Weise die „Detection and Response“-Fähigkeiten über E-Mail, Netzwerke, Endpunkte, Server und Cloud-Workloads hinweg ausbauen. Unternehmen sollen so einen besseren und umfassenden

Überblick über ihren Sicherheitsstatus erhalten. Gleichzeitig können sie kleinere Vorfälle aus verschiedenen Sicherheitssilos miteinander in Verbindung bringen, um komplexe Angriffe zu erkennen, die sonst unentdeckt bleiben würden. Umfragen zufolge setzen laut Trend Micro 55 Prozent der Unternehmen mehr als 25 ver-

schiedene Cybersecurity-Techniken ein. Dennoch schaffen es die beständig zunehmenden Cyberangriffe weiterhin, bestehende Sicherheitsmechanismen zu umgehen. IT-Security-Teams erhalten täglich bis zu 10.000 Hinweise auf Vorfälle. Dies erzeugt eine hohe Belastung und macht ihre Arbeit ineffizient. Laut dem Verizon 2018 Data Breach In-

vestigations Report beträgt die durchschnittlich benötigte Zeit bis zur Identifikation eines Datenabflusses 197 Tage. Bis das Abfließen gestoppt ist, kann es weitere 69 Tage dauern. Cyberkriminelle haben im Schnitt also fast neun Monate Zeit. jos

Info: Trend Micro
Tel.: 0811/88990-700
Web: www.trendmicro.de