

Negative Laufzeit – gibt's die wirklich?

Manfred Zollner

Im "Bode-Diagramm" stellt man das Übertragungsverhalten eines Systems durch den Betrags- und den Phasenfrequenzgang dar. Aus der Phase können die Phasen- und die Gruppenlaufzeit abgeleitet werden. Bei einer Reihe von Signalen beschreiben diese beiden Größen sehr anschaulich die vom System hervorgerufene zeitliche Verzögerung. Dies gelingt aber nicht bei allen Signalen gleich gut, und insbesondere die negative Laufzeit bereitet Verständnisprobleme. Anhand der folgenden Beispiele werden die o.a. Systemgrößen vorgestellt und in ihrem Bezug zur Zeitfunktion verdeutlicht.

Die Nachrichtentechnik beschreibt die Übertragung eines Signals von einer **Quelle** (Sender) durch einen **Kanal** zu einer **Senke** (Empfänger). Der Kanal kann eine Funkstrecke sein, ein Kabel, ein Verstärker, ein Lautsprecher, letztlich alles, was zwischen Quelle und Senke die Signalübertragung ermöglicht. Stehen die Veränderungen, die das Signal bei der Übertragung erfährt, im Vordergrund, übernimmt die Systemtheorie: Der Kanal wird zum **System**, das Signale aufeinander abbildet. Die signalbeschreibenden Größen sind Zeitfunktion ($\underline{x}(t)$, $\underline{y}(t)$) und zugehöriges Spektrum ($\underline{X}(j\omega)$, $\underline{Y}(j\omega)$), die systembeschreibenden Größen sind Impulsantwort $\underline{h}(t)$ und Übertragungsfunktion $\underline{H}(j\omega)$. Die verschiedenen Teilgebiete der Systemtheorie sind umfangreich, jedoch in der Literatur gut dokumentiert [z.B. 1]. Im Folgenden soll nur ein kleiner Ausschnitt daraus beleuchtet werden: Die **Verzögerung**, die ein Signal auf seinem Weg von der Quelle zur Senke erfährt.

Im allgemeinen Fall sind Signal- und Systemgrößen komplexwertig, was durch Unterstreichen dargestellt wird, z.B. $\underline{x}(t) = \text{Re}\{\underline{x}(t)\} + j\text{Im}\{\underline{x}(t)\}$. Mit Re wird der Realteil bezeichnet, mit Im der Imaginärteil, $j = \sqrt{-1}$ ist die imaginäre Einheit, in der Mathematik auch i genannt. Die Beschränkung auf reelle Zeitfunktionen (z.B. $x(t)$) stellt zunächst keine Beeinträchtigung dar, es muss aber betont werden, dass die zugehörigen Spektralfunktionen trotzdem komplexwertig sein können. In **Abb. 1** ist ein reeller Rechteckimpuls dargestellt. Der Begriff **Puls** sollte hierfür nicht verwendet werden, ein Puls ist ein periodisches Signal. Das rechte Bild zeigt den gleichen Impuls, aber zeitlich um 2 ms verschoben. Nach *rechts* verschoben, d.h. verzögert. Das Leben lehrt, dass Wirkungen nicht vor der Ursache kommen können, dass Signalabbildungen in der Realität deshalb **kausal** sein müssen. Und sie können auch niemals so

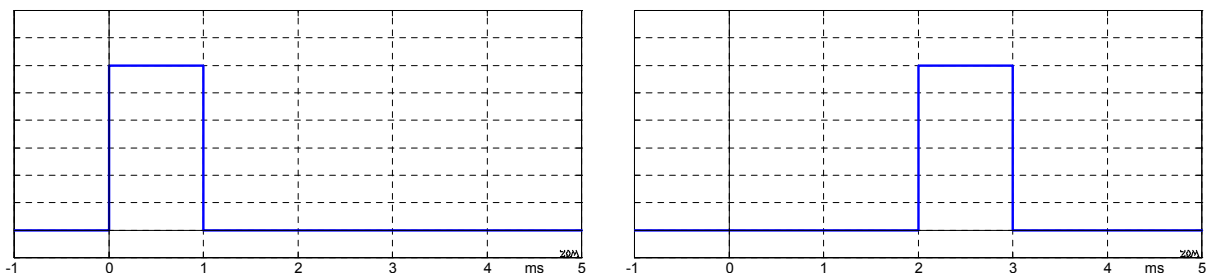


Abb. 1: Rechteckimpuls am Eingang (links) und Ausgang (rechts) eines Systems, Laufzeit = 2 ms.

Die restlichen Seiten sind als PDF downloadbar: www.gitec-forum.de

ideal sein wie in Abb. 1: Eine Formänderung muss der Impuls zusätzlich zu seiner Verschiebung auch erfahren, denn sonst wäre das System unendlich breitbandig, was reale Systeme nicht sein können. **Abb. 2** zeigt hierzu als Beispiel die Impulsabbildung durch ein Tiefpassfilter: Im linken Bild den Rechteckimpuls als Filter-Eingangssignal, im rechten Bild das zugehörige Filter-Ausgangssignal. Wie stark, also um wie viele Millisekunden verzögert der Tiefpass, wie groß ist die Laufzeit (Verzögerungszeit)? Der Impulsanfang wird wieder um 2 ms verzögert, aber der Ausgangsimpuls "hängt irgendwie nach rechts", sein Schwerpunkt wird um mehr als 2 ms verzögert. Aber ist die Schwerpunktverschiebung das richtige Kriterium? Oder präziser: wie definiert die Systemtheorie die in Systemen auftretende Verzögerung?

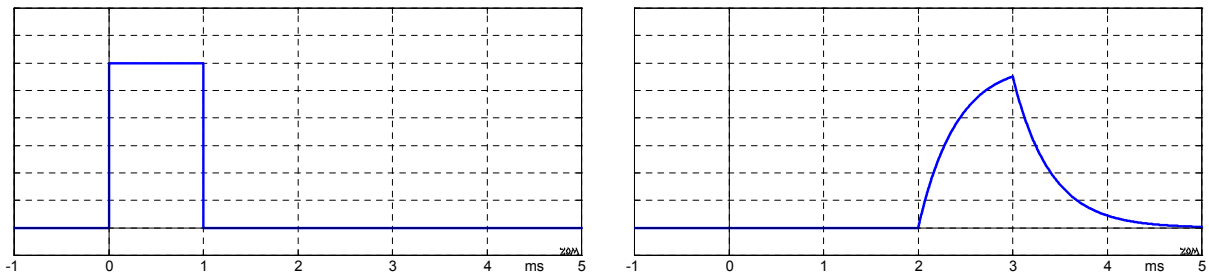


Abb. 2: Impulsabbildung durch Verzögerung und Bandbegrenzung.

Durchaus so, wie oben angegeben, aber nicht nur. Zunächst zum einfachen Fall: Sofern das System nur eine Verzögerung bewirkt (was nur bei idealisierten Systemen möglich ist), kann die Verzögerung als Koordinatentransformation (Verschiebung, Translation) mit der Laufzeit τ beschrieben werden: $y(t) = x(t - \tau)$. Hier sind x und y formgleich, aber horizontal gegeneinander verschoben, wie in Abb. 1. Findet zusätzlich eine Formänderung statt, wird ein Umweg über das **Spektrum** erforderlich. Zum System-Eingangssignal $x(t)$ gehört das Spektrum $X(j\omega)$, es kann aus $x(t)$ z.B. mittels Fourier-Transformation ermittelt werden. Entsprechendes gilt beim Ausgangssignal. Der Quotient aus den beiden Spektren ist die Übertragungsfunktion $H(j\omega) = Y(j\omega)/X(j\omega)$. Auch bei reellen Signalen ist die Übertragungsfunktion so gut wie immer komplex, sie kann in einen frequenzabhängigen **Betrag** $H(\omega)$ und eine frequenzabhängige **Phase** $\varphi(\omega)$ zerlegt werden. Aus dem spektralen Differential der Phase ergibt sich – wie später noch genauer dargelegt wird – die **Gruppenlaufzeit**, und auch die ist frequenzabhängig. Das bedeutet: Treten bei der Impulsübertragung Formänderungen auf, ist zur Beschreibung der Verzögerung ein einziger Zahlenwert nicht mehr ausreichend; vielmehr muss eine von der Frequenz abhängige Verzögerungsfunktion angegeben werden.

Man könnte versucht sein, die Frequenzachse als Summe unendlich vieler Frequenzpunkte aufzufassen, also das Kontinuum Frequenz zu diskretisieren. Mathematisch nicht falsch, aber für die folgenden Überlegungen nicht ganz ungefährlich. Weil eine Sinusfunktion (Bandbreite null) sich in ihrem Verlauf nicht ändern kann. Soll sich z.B. die Amplitude ändern, ist eine von null verschiedene Bandbreite erforderlich. Als Beispiel betrachten wir einen kosinusförmig amplitudenmodulierten Kosinuston (**Abb. 3**). Die blaue Kosinusschwingung ändert

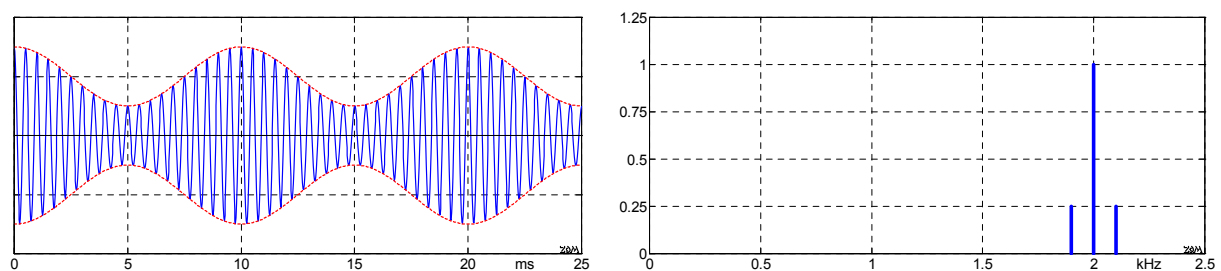


Abb. 3: Mit Kosinus-Hüllkurve amplitudenmodulierter Kosinuston, rechts das zugehörige Betragsspektrum.

ihre Amplitude über der Zeit, als Hilfslinie ist rot die **Hüllkurve** eingezeichnet. Aus den zeitlichen Abständen der Nulldurchgänge lässt sich die Periodendauer und als deren Kehrwert die Frequenz (2 kHz) dieser blauen Kosinusschwingung ermitteln, die im Folgenden "Träger" genannt wird. Der Begriff kommt aus der Funktechnik: Diese Schwingung "trägt" die Nachricht (die Hüllkurve) durch die Luft zur Empfangsantenne. Die Hüllkurve existiert im Signal nicht, sie ist eine gedachte Linie. Eine Frequenz kann ihr trotzdem zugeordnet werden, im Beispiel ist das 100 Hz. Die Hüllkurve wird auch Modulation genannt, oder Nutzsignal, oder einfach nur NF (Niederfrequenz), der Träger heißt manchmal dementsprechend HF (Hochfrequenz). Formal entsteht dieses Modulationssignal durch Multiplikation des Trägers mit einem durch Offset verschobenen NF-Signal:

$$x(t) = [1 + 0.5 \cos(\omega_N \cdot t)] \cdot \cos(\omega_H \cdot t); \quad \text{Amplitudenmodulation = AM}$$

Das Produkt der beiden Winkelfunktionen lässt sich in eine Summe umformen:

$$x(t) = \cos(\omega_H \cdot t) + 0.25 \cos(\omega_D \cdot t) + 0.25 \cos(\omega_S \cdot t); \quad \text{Summenformel}$$

Hierbei steht H jeweils für die Hochfrequenz, N für die Niederfrequenz, D für die Differenz der beiden Frequenzen und S für deren Summe. Das AM-Signal lässt sich also auf zwei verschiedene Arten erzeugen: Als Produkt dreier Signale (Offset, f_N , f_H), oder als Summe dreier anderer Signale (f_D , f_H , f_S). In einem **Spektrum** (wie z.B. in Abb. 1) wird immer die Summen-darstellung visualisiert, niemals die Produktdarstellung. Das o.a. AM-Signal soll nun durch ein System übertragen werden, in dem die Phasen der Teilschwingungen verändert werden; die Amplituden (Beträge) bleiben unverändert. Hierbei gibt es einen Spezialfall: wenn nämlich jedes Signal – unabhängig von seiner Frequenz – um dieselbe Zeit τ verzögert wird. Die hierfür erforderliche Phasenverschiebung kann dann leicht berechnet werden:

$$x(t) = \cos(\omega \cdot (t - \tau)) = \cos(\omega t + \varphi); \quad \varphi = -\omega \cdot \tau;$$

Bei einer frequenzunabhängigen Verzögerung um τ ist die Phase proportional zur Frequenz. **Abb. 4** zeigt so einen Fall: Die aus Abb. 1 bekannte AM, im rechten Bild um τ verzögert. In dieser Abbildung wurden Träger und Hüllkurve um dieselbe Zeit τ verzögert, weshalb die Form der AM erhalten bleibt – sie wandert nur als Ganzes nach rechts. Doch ist dies ein Sonderfall, auf einen Freiheitsgrad (nämlich τ) reduziert. In der o.a. Summenformel gibt es aber drei Einzelschwingungen, und jede könnte eine eigene, von den anderen unabhängige Phasenänderung erfahren. Setzt man tatsächlich drei verschiedene Phasenverschiebungen an, und versucht, die Summenformel wieder in die Produktformel zurückzutransformieren, so gelingt dies nur, wenn $(\varphi_D + \varphi_S)/2 = \varphi_H$ eingehalten wird: die Phasenänderung des Trägers muss dem Mittelwert der beiden anderen Phasenänderungen entsprechen. Die Herleitung dieser Bedingung gelingt sehr leicht mittels der komplexen Signaldarstellung, mehr dazu am Ende.

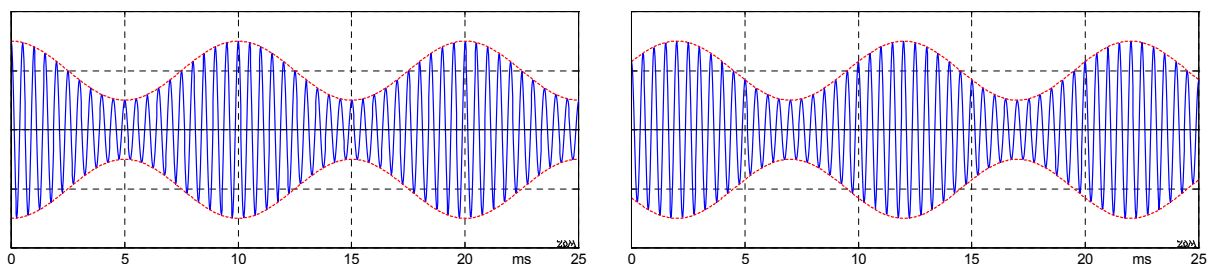


Abb. 4: Mit Kosinus-Hüllkurve amplitudenmodulierter Kosinuston, rechts um 2 ms verzögert.

Mit dieser Bedingung (Rücktransformierbarkeit) bleiben zwei Freiheitsgrade für die Phase:

$$x(t) = \cos(\omega_H \cdot t + \varphi_H) + 0.25 \cos(\omega_D \cdot t + \varphi_H - \Phi) + 0.25 \cos(\omega_S \cdot t + \varphi_H + \Phi);$$

Für das Phasenspektrum bedeutet dies, dass alle drei Phasen auf einer Geraden liegen müssen. Diese kann durch den Ursprung gehen, muss aber nicht. **Abb. 5** zeigt hierzu zwei Beispiele: Im linken Bild geht die rote Gerade durch den Ursprung, im rechten Bild nicht. Die beiden Freiheitsgrade sind: die Phase des Trägers (φ_H), und die Abweichung von diesem Wert ($\pm\Phi$).

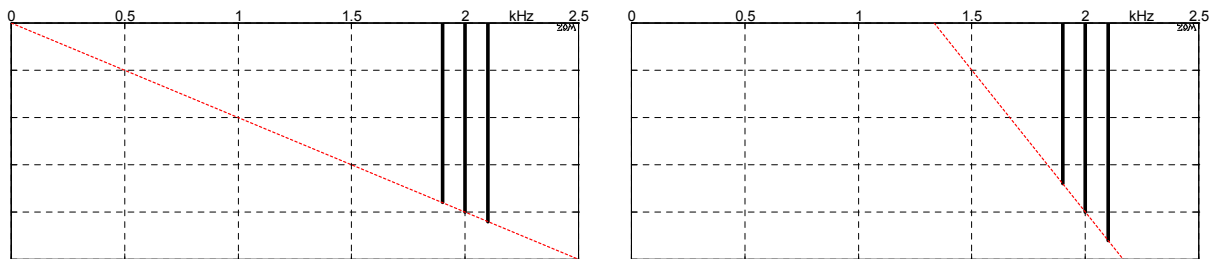


Abb. 5: Phasenspektren. Die Phasen der drei Teilschwingungen liegen jeweils auf einer (roten) Geraden.

Die Rücktransformation der o.a. Summenformel führt mit φ_H und Φ auf die Produktformel:

$$x(t) = [1 + 0.5 \cos(\omega_N \cdot t + \Phi)] \cdot \cos(\omega_H \cdot t + \varphi_H)$$

Hier sieht man sehr schön die beiden Freiheitsgrade: Mit φ_H kann die Trägerphase verschoben werden, mit Φ die Hüllkurvenphase – unabhängig voneinander. Diese beiden Verschiebungen (Verzögerungen) sind in **Abb. 6** dargestellt. Im linken Bild bleibt die Hüllkurve unverändert $HK = [1 + 0.5 \cos(\omega_N \cdot t)]$, aber die Trägerphase (φ_H) wird verschoben. Im rechten Bild bleibt die Trägerphase unverändert $\cos(\omega_H \cdot t)$, hier wird die Hüllkurve verzögert.

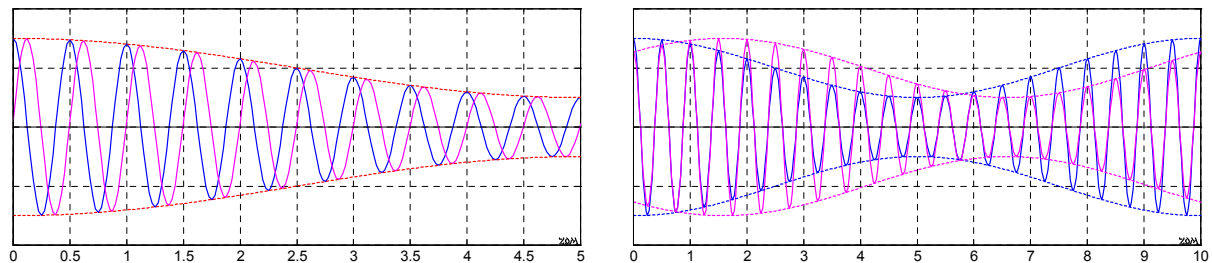


Abb. 6: Verschiedene Verzögerungen: Gemeinsame Hüllkurve (links), gemeinsame Trägerphase (rechts).

Die Verzögerung des Trägers wird **Phasenlaufzeit** τ_P genannt, die Verzögerung der Hüllkurve **Gruppenlaufzeit** τ_G . Die Phasenlaufzeit ergibt sich zu $\tau_P = -\varphi_H / \omega_H$, die Gruppenlaufzeit zu $\tau_G = -d\varphi/d\omega$, also als negatives spektrales Differential der Phase, bzw. als negative Steigung* der in Abb. 5 rot eingezeichneten Geraden. In der Nachrichtentechnik ist die Gruppenlaufzeit die wichtigere der beiden, denn in der Hüllkurve ist die zu übertragende Nachricht codiert – zumindest bei der hier betrachteten Amplitudenmodulation. Diese Hüllkurve muss keineswegs nur einer Kosinusfunktion folgen. In **Abb. 7** sind zwei Signale mit etwas komplizierteren Spektren dargestellt. Im ersten Beispiel besteht die Hüllkurve aus zwei überlagerten Kosinusschwingungen (100 Hz und 200 Hz), im zweiten Beispiel, einem Gaußton, ist das Spektrum kontinuierlich (siehe [2]). Doch auch bei diesen Signalen lassen sich die o.a. Laufzeiten berechnen, auch diese Signale können als AM in der Summen- oder Produktform dargestellt werden.

* Ganz exakt: Die Phase wird nach ω differenziert, nicht nach f .

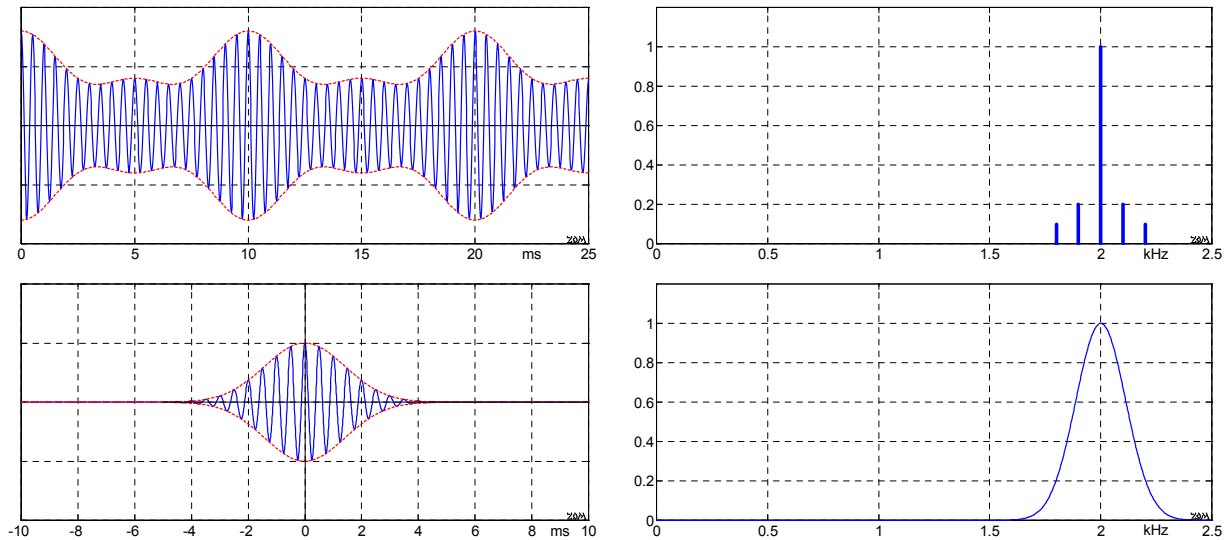


Abb. 7: Zeitfunktion und Betragsspektrum einer nichtsinusförmigen Amplitudenmodulation. Das Spektrum in der ersten Bildzeile ist diskret, das in der zweiten Bildzeile ist kontinuierlich.

Die Unterschiede zwischen Phasen- und Gruppenlaufzeit sollen nochmals anhand von **Abb. 8** verdeutlicht werden. Die dargestellten Phasenfrequenzgänge sind die des übertragenden Systems, durch sie wird das Signal von Abb. 3 übertragen (abgebildet). Im ersten Beispiel **a**) ist die Phase frequenzunabhängig, $\tau_G = 0$, $\tau_P = (3\pi/4)/(2\pi \cdot 2\text{kHz}) = 187.5\mu\text{s}$; die Hüllkurve bleibt unverändert, der Träger wird um τ_P verzögert. Im zweiten Beispiel **b**) erfährt der Träger keine Phasenverschiebung (= keine Verzögerung), jedoch wird die Hüllkurve um 1.3 ms verzögert. Bei dritten Beispiel **c**) geht die rote Gerade durch den Ursprung, Phasen- und Gruppenlaufzeit sind gleich ($187.5\mu\text{s}$). Im vierten Beispiel **d**) ist die Phasenlaufzeit genau so groß wie bei c), die Gruppenlaufzeit ist jedoch größer (steiler verlaufende Gerade) $\tau_G = 750\mu\text{s}$.

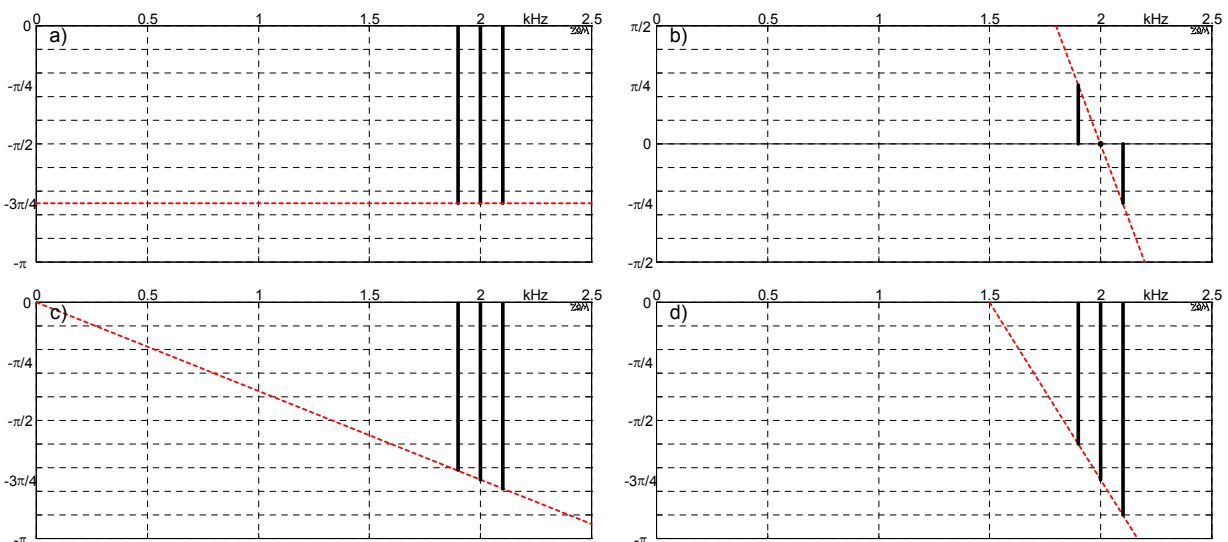


Abb. 8: Beispiele für vier unterschiedliche Phasenfrequenzgänge.

Bei allen vier Beispielen aus Abb. 8 liegen die drei Phasen auf einer (roten) Geraden, deshalb bleibt die Form der Hüllkurve erhalten – sie wird lediglich verschoben. Was passiert, wenn von dieser Bedingung abgegangen wird, zeigt **Abb. 9**. Hierbei wird nur die obere Seitenlinie (2.1 kHz) um $-\pi$ phasenverschoben, wodurch die Amplitudenmodulation fast ganz verschwindet. Es entsteht sogar eine leichte Frequenzmodulation, die im Bild aber nicht zu erkennen ist.

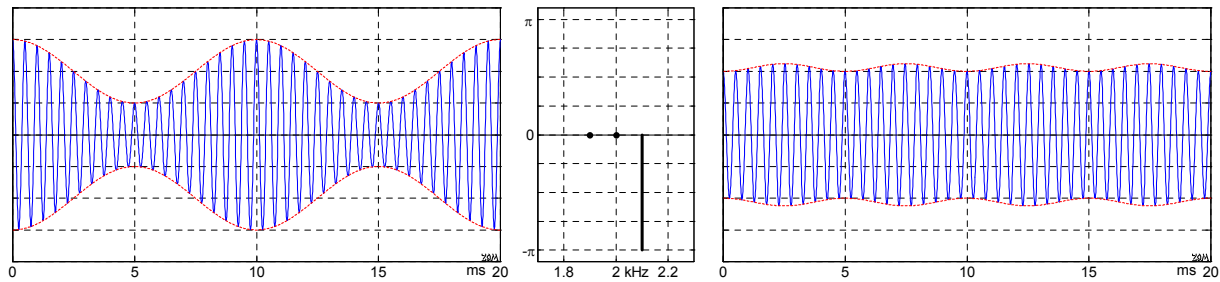


Abb. 9: Abbildung der AM (links) durch eine formändernde Phasenverschiebung.

Abb. 9 macht deutlich, dass die Aussagekraft der Phasen- und Gruppenlaufzeit doch etwas begrenzt ist. Wenn das zu übertragende Signal als amplitudenmodulierte Trägerschwingung darstellbar ist, und wenn außerdem die Phasenfunktion eine Gerade ist, kann aus der Phasenlaufzeit auf die Trägerverschiebung geschlossen werden, und aus der Gruppenlaufzeit auf die Hüllkurvenverschiebung. Die Darstellung als amplitudenmodulierte Trägerschwingung ist immer dann möglich, wenn das Betragsspektrum achsensymmetrisch um die Trägerfrequenz verläuft, und das Phasenspektrum punktsymmetrisch um die Trägerfrequenz. Bei den Signalen in Abb. 7 war das der Fall, bei dem Rechteckimpuls von Abb. 1 nicht. Wie wirken sich bei ihm Phasenverschiebungen aus? Hier und im Folgenden ist eine Unterscheidung nötig zwischen Phasenverschiebungen, die als Folge von Amplitudenänderungen (zwangsläufig) auftreten, und reinen Phasenverschiebungen (ohne Amplitudenänderung).

Amplitudenänderungen können **minimalphasig** oder **allpasshaltig** erfolgen [3]. Bei minimalphasigen Filterungen muss sich auch die Phase ändern, sobald der Betragsfrequenzgang von 0 dB abweicht. **Allpässe** ermöglichen hingegen eine Phasenänderung, ohne dass sich der Betrag ändert. Der Rechteckimpuls in **Abb. 10** wird minimalphasig tiefpassgefiltert. Zu hohen Frequenzen hin konvergieren beide Laufzeitfrequenzgänge gegen null. Der hochfrequente Anteil des Rechteckimpulses ist der Knick, und der wird unverzüglich übertragen. Tieffrequent konvergieren beide Laufzeiten gegen 0.4 ms. Kein offensichtlicher Widerspruch, aber es hätten auch nur 0.3 ms sein können – die Verformung des Impulses stellt kein eindeutiges Kriterium mehr zur Laufzeitdefinition dar. Und nicht zu vergessen: Der Tiefpass verändert ja auch das Betragsspektrum, er schwächt die hochfrequenten Anteile ab, was auch zur Impulsverformung beiträgt. Bereits bei diesem Beispiel sieht man: So klar, wie bei AM-Signalen, lassen sich die beiden Laufzeiten nicht mehr der Zeitfunktion zuordnen.

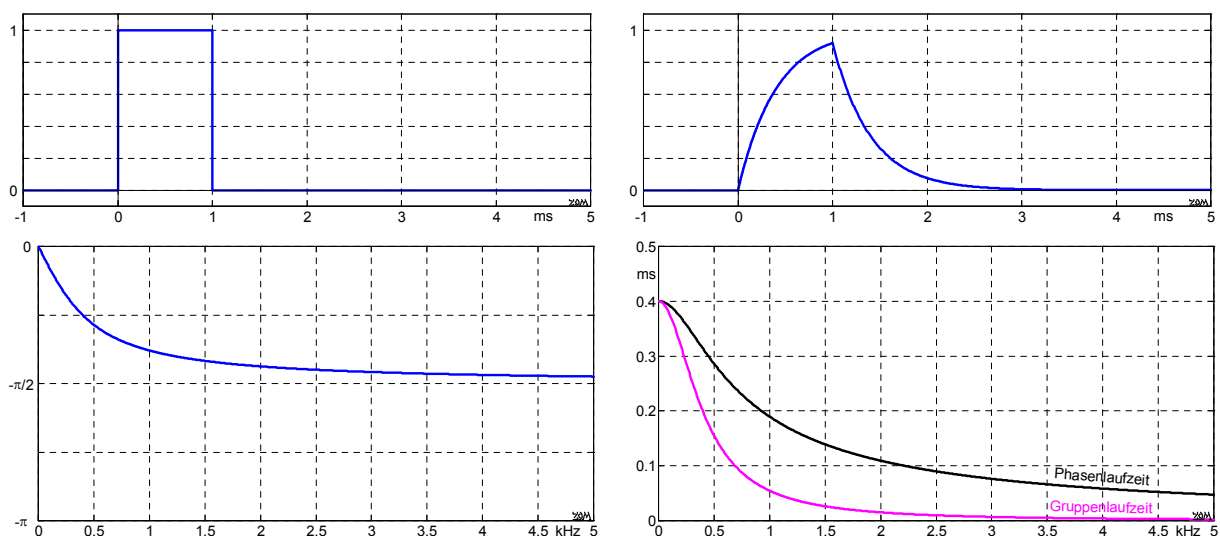


Abb. 10: Impuls, Tiefpass 1. Ordnung; Phasenfrequenzgang (l. u.), Gruppen- und Phasenlaufzeit (r. u.).

Nun zu einem **Allpass 1. Ordnung (Abb. 11)**. Der Tiefpass hatte den Betrag verändert – der Allpass dreht nur die Phase, und lässt den Betrag, wie er ist. Abgebildet wird wieder der bekannte Rechteckimpuls, die Unterschiede zu Abb. 10 sind erheblich. Dabei haben sich Phase, Gruppen- und Phasenlaufzeit lediglich verdoppelt – aber zusätzlich ist die vom Tiefpass verursachte Höhendämpfung weggefallen. Die hohen Frequenzanteile (der Sprung) werden gegenphasig wiedergegeben (Phasendrehung auf $-\pi$). Deshalb springt die Allpassantwort von 0 auf -1 , und nicht auf $+1$ wie beim Rechteckimpuls. Die tiefen Frequenzanteile werden demgegenüber (als Grenzfall) ohne Phasendrehung abgebildet – deshalb wechselt die Allpassantwort bei 0.25 ms das Vorzeichen.

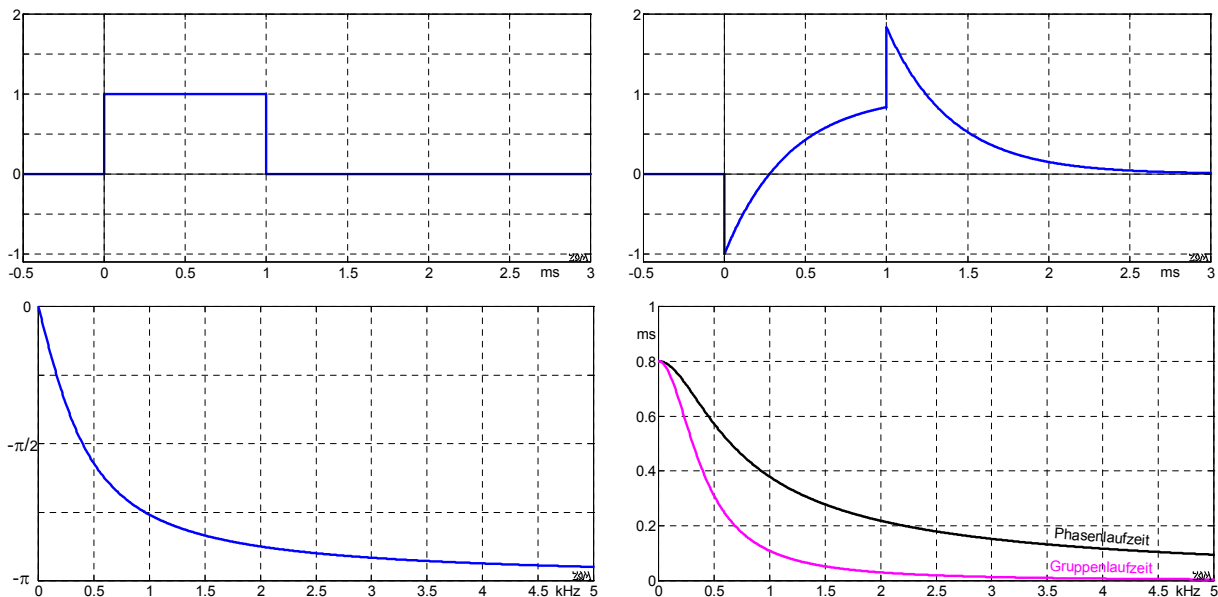


Abb. 11: Durch einen Allpass 1. Ordnung hervorgerufene Impulsverformung. Phase, Gruppen- und Phasenlaufzeit sind genau doppelt so groß wie in Abb. 10.

Nun sollen zwei in Kette geschaltete Systeme untersucht werden, eine Anordnung, die in ähnlicher Weise als Pre-/Deemphasis hilft, den Rauschabstand zu verbessern. Das erste System erzeugt eine Höhenanhebung, das zweite eine hierzu inverse Höhenabsenkung (**Abb. 12**). Die Übertragungsfunktion des zweiten Systems ist invers (reziprok) zu der des ersten $H_2 = 1/H_1$, am Ausgang des zweiten Systems liegt deshalb exakt dasselbe Signal an wie am Eingang des ersten Systems. Bei Kettenschaltung zweier Systeme sind sowohl die Verstärkungsmaße (die dB-Werte), als auch die Phasen zu addieren; ihre jeweilige Summe ergibt sich für jede Frequenz zu null. In der Schaltungstechnik heißt das erste System "Hochpass mit Gegenhalt", das zweite "Tiefpass mit Gegenhalt" [2].

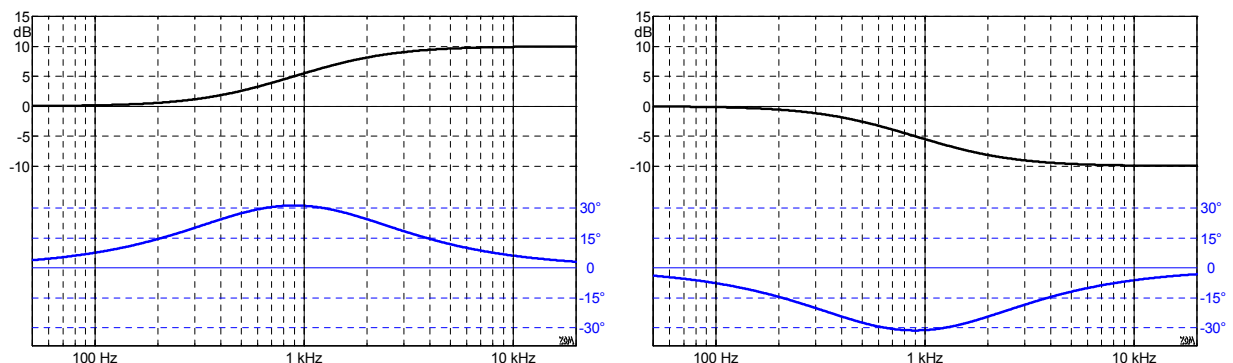


Abb. 12: Kettenschaltung zweier Systeme: Höhenanhebung (links), hierzu inverse Höhenabsenkung (rechts).

Betrags- und Phasenfrequenzgang in Abb. 12 geben keinen Anlass zu Kritik, doch bei den Gruppenlaufzeitfrequenzgängen (**Abb. 13**) entsteht die Frage: **Negative Laufzeit, gibt's die?** Das würde ja bedeuten, dass am Ausgang ein Signal erscheint, noch bevor das Eingangssignal anliegt! Das wäre ein **akausales System**, bei dem die Wirkung vor der Ursache eintritt?

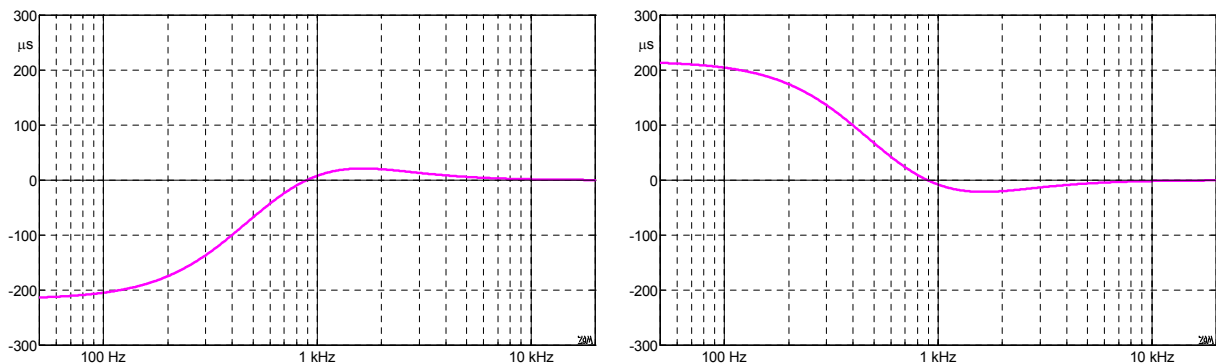


Abb. 13: Gruppenlaufzeitfrequenzgänge der beiden Systeme aus Abb. 12.

Aber wie sonst soll das sein? Das erste System ändert den Betrag des Signals, minimalphasig, also muss sich die Phase ändern, nach festen Regeln, mit einer Gruppenlaufzeit. Wäre die rein positiv, müsste das zweite System eine rein negative Gruppenlaufzeit haben, damit sich in summa wieder null ergibt. Doch so unrealistisch ist eine negative Gruppenlaufzeit gar nicht: hiermit wird ja der **eingeschwungene Zustand** bezeichnet, und bis der erreicht ist, bleibt genügend Zeit, dass das Signal "etwas nach links rutschen kann". In **Abb. 14** ist die Filterung eines Burstsignals (schwarz) dargestellt [3]*. Durch die negative Gruppenlaufzeit rutscht der Schwerpunkt der Bursts nach links, aber nicht vor den Beginn (0 ms) – auch diese Filterung ist (trotz negativer Gruppenlaufzeit) kausal.

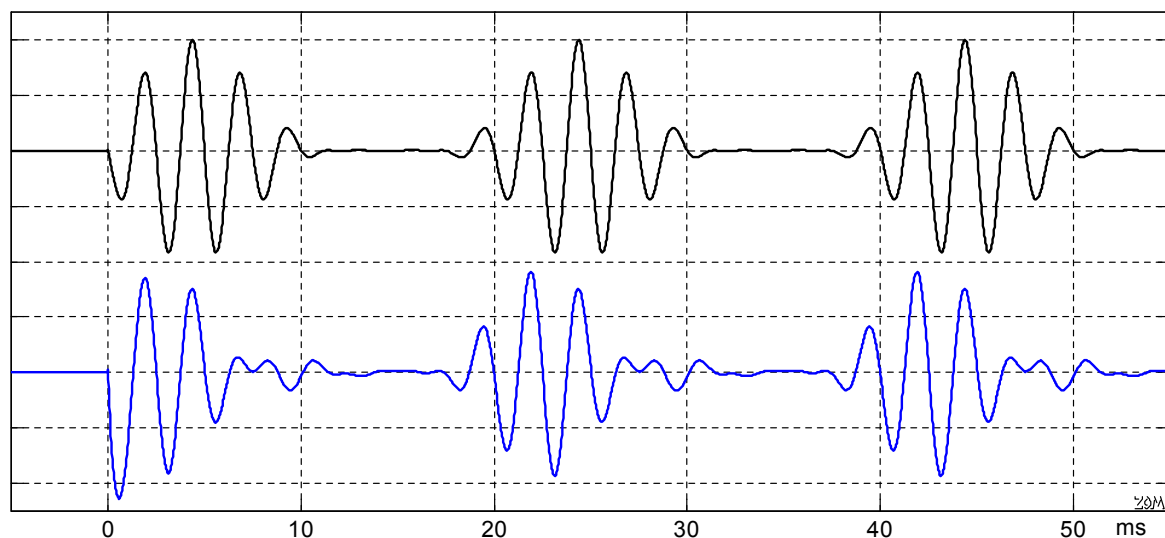


Abb. 14: Negative Gruppenlaufzeit; Generatorsignal = schwarz, Filterausgang = blau.

Gruppen- und Phasenlaufzeit gehören zum Standardinstrumentarium der Nachrichtentechnik, ihre Interpretation ist aber nicht ganz so einfach wie die des Betragsfrequenzganges. Sie sind als (eigentlich redundante) Ergänzung zum Phasenfrequenzgang aufzufassen, und können die visuelle Auswertung unterstützen. Entsprechende Erfahrung vorausgesetzt.

* Weil bei diesem Beispiel auch der Betrag verändert wird, ändert sich die Kurvenform etwas.

Anhang: Komplexe Signaldarstellung

In der zweiseitigen Spektraldarstellung hat ein reeller Kosinuston eine positive und eine negative Frequenz [2]: $\cos(\omega t) = 0.5 \cdot (\exp(j\omega t) + \exp(-j\omega t))$. Die Multiplikation zweier Kosinusschwingungen (NF und HF) führt zu einer Addition der Exponenten, d.h. zu einer spektralen Verschiebung (siehe auch Faltung im Frequenzbereich). Reelle Zeitsignale haben ein gerades Betrags- und ein ungerades Phasen-Spektrum (Zuordnungssatz der Fourier-Transformation), diese Symmetrieeigenschaften bleiben beim Verschieben (Falten) erhalten. Daraus folgt im Umkehrschluss: Wenn Signale (wie bei der AM) als Produkt zweier reeller Funktionen dargestellt werden sollen, muss ihr Betragsspektrum achsensymmetrisch (gerade) zu einer Trägerfrequenz sein, und ihre Phasenspektrum punktsymmetrisch (ungerade). Für eine ausführliche Darstellung wird auf [2] und [4] verwiesen.

Literatur

- [1] Marko H.: Methoden der Systemtheorie, Springer 1986,
- [2] Zollner M.: Signalverarbeitung, Bibliothek der OTH Regensburg, ausleihbar.
- [3] Zollner M.: Equalizer und Allpässe, www.gitarrenphysik.de
- [4] Zollner M.: Frequenzanalyse, Bibliothek der OTH Regensburg, ausleihbar.

Fachartikel der GITEC Knowledge Base:	
1 Gitarren-Lautsprecher	11 Schaltungsvarianten für das Reguitaring
2 Studio-Lautsprecher	12 Verzerrungen: gerade oder ungerade?
3 Welche ECC83 darf's denn sein?	13 Die Basswiedergabe beim Studio-Monitor
4 Reamping and Reguitaring	14 Vom Sinn und Unsinn der CSD-Wasserfälle
5 Gitterstrom bei Trioden	15 Artefakte bei Wasserfall-Spektrogrammen
6 Der Verzerrer	16 Equalizer und Allpässe, Teil 1 – 3
7 Der Range-Master rauscht	17 Studio- und Messmikrofone, Teil 1 – 5
8 Raumakustik	18 Die Dummy-Load als Lautsprecher-Ersatz
9 Saitenalterung	19 Nichtlineare Modelle
10 Lautsprecherkabel	20 Wie misst man Elkos?
	21 Der Lautsprecher-Phasengang