

**Vorlesung „Einführung in die Wirtschaftspolitik“ für Lehramtsstudiengänge
und Politikwissenschaften, WS 2016/2017**

Teil 2: Makroökonomische Theorie

1. Einleitung

Anders als die mikroökonomische Theorie, die sich mit dem Verhalten einzelner Wirtschaftssubjekte (Haushalte und Unternehmen) bzw. der Funktionsfähigkeit einzelner Märkte beschäftigt, konzentriert sich die makroökonomische Theorie auf gesamtwirtschaftliche Prozesse – also unter Vernachlässigung dessen, was sich auf der mikroökonomischen Ebene abspielt. Makroökonomische Zusammenhänge beruhen zwar auf dem Verhalten einzelner Wirtschaftssubjekte (so dass es keine Widersprüche zwischen beiden Betrachtungsweisen gibt oder geben sollte), führen aber zu Ergebnissen, die über eine bloße Aggregation einzelwirtschaftlicher Anpassungsprozesse hinausgehen.

Die klassische (oder auch: neoklassische) makroökonomische Theorie beruht in ihrem Kern auf einer Übertragung mikroökonomischer Verhaltensweisen auf die gesamtwirtschaftliche Ebene: Durch Zusammenfassung aller einzelnen Gütermärkte (und der dort relevanten Angebots- und Nachfrage) ergibt sich ein aggregierter Gütermarkt (mit dem typischen fallenden Verlauf der Nachfrage bzw. dem typischen steigenden Verlauf des Angebots), durch Zusammenfassung aller Faktormärkte (vereinfacht: aller Arbeitsmärkte) ergibt sich ein aggregierter Faktormarkt (Arbeitsmarkt). Da die Arbeitnehmer nur so viel Arbeit anbieten, dass sie damit ihre Konsumwünsche erfüllen können, folgt, dass aggregierter Arbeitsmarkt und aggregierter Gütermarkt stets in einem Gleichgewicht sind; dieses wird durch korrespondierende Anpassungen von Preisniveau und Lohnniveau erreicht. Temporäre Ungleichgewichte (z.B. aufgrund von unerwarteten Nachfrageschwankungen, Präferenzänderungen oder ähnlichem) werden also durch Preis- und Lohnanpassungen abgebaut. Für „einfache“ Analysezwecke mag dies reichen; es entspricht aber offenkundig nicht der Realität, in der es ja auch länger andauernde Ungleichgewichtssituationen (wie z.B. lang anhaltender Arbeitslosigkeit)

gibt. Diese Diskrepanz zwischen Empirie und theoretischer Vorhersage führte zur Entwicklung einer eigenständigen makroökonomischen Betrachtung, die eng mit dem Namen von John Maynard Keynes (1883-1946) verknüpft ist.

2. Einführung: Die Rolle wirtschaftspolitischer Akteure (Staat und Zentralbank)

Eine wesentliche Erweiterung der „modernen“ makroökonomischen Theorie gegenüber der mikroökonomischen Herangehensweise ist es, dass nunmehr weitere Akteure zu berücksichtigen sind, nämlich der Staat (=Fiskalpolitik) und die Zentralbank (=Geldpolitik).

Der Staat handelt der nicht notwendigerweise nach ökonomischen Optimalitätskriterien (wie Nutzen- oder Gewinnmaximierung), sondern kann mit seinen Aktivitäten (auch) wirtschaftspolitische Ziele verfolgen. Wesentliches Instrument sind dabei die Ausgaben- und die Einnahmenpolitik. Bei den Ausgaben wird dabei typischerweise nach drei Zielsetzungen unterschieden, nämlich die Allokation (=Beeinflussung der Marktergebnisse im Falle von Marktversagen oder einer unerwünschten Ressourcenverteilung, z.B. durch Subventionen; aber auch: Güterverkäufe des Staates zur Bereitstellung öffentlicher Güter); die Stabilisierung (durch Kompensation von konjunkturellen Schwankungen durch Variation von Ausgaben) und die Distribution (Korrektur von Marktergebnissen unter sozialpolitischen Gesichtspunkten, beispielsweise durch Zahlung von Sozialtransfers). Bei den Einnahmen können ebenfalls diese Zielsetzungen verfolgt werden (z.B. durch Gewährung von Steuerermäßigungen (=Allokation), durch temporäre Steuersatzsenkungen/-erhöhungen (=Stabilisierung) oder durch Festlegung von unterschiedlich hohen Steuersätzen zur Korrektur der primären Einkommensverteilung (=Distribution); darüber hinaus dienen die Einnahmen aber auch der Erwirtschaftung der benötigten Mittel für die Finanzierung der als notwendig oder wünschenswert angesehenen Ausgaben. Neben Steuereinnahmen kommt dabei auch eine Defizitfinanzierung öffentlicher Ausgaben in Betracht.

Zum zweiten muss bei einer gesamtwirtschaftlichen Betrachtung auch die Zentralbank berücksichtigt werden, die das Geld emittiert, das für Zahlungszwecke verwendet wird; die Geldschöpfung kann ebenfalls dazu verwendet werden, wirtschaftspolitische Ziele zu erreichen, so zum Beispiel indem eine Ausweitung der Geldmenge die gesamtwirtschaftliche Nachfrage erhöhen und damit unter bestimmten Bedingungen

expansive realwirtschaftliche Impulse auslösen kann (bei der mikroökonomischen Betrachtung spielte „Geld“ keine Rolle, weil die dort im Vordergrund stehenden Überlegungen zum Nutzen- bzw. Gewinnmaximierungskalkül und damit die abgeleiteten Marktprozesse immer nur auf relativen Preisen einzelner Güter (Preis eines Gutes in Werteinheiten eines anderen Gutes) beruhten).

Geld kann alles sein, was bestimmte Kriterien erfüllt: Es muss allgemein als werthaltig angesehen werden, und es darf nicht von jedermann in beliebiger Höhe emittiert werden. In früherer Zeit erfüllten beispielsweise auch Muscheln oder Perlen (auf irgendwelchen Südseeinseln) diesen Zweck; nach dem 2. Weltkrieg auch Zigaretten. Geld muss auch nicht notwendigerweise von dem Staat ausgegeben werden, in dem es umläuft – beispielsweise gab es in der DDR auch Geldkreisläufe von D-Mark, in vielen Entwicklungsländern kann man heute noch mit Dollar oder Euro bezahlen.

Geld erfüllt in wirtschaftlicher Betrachtung mehrere Funktionen: Am bedeutsamsten ist dabei die „Zahlungsmittelfunktion“: Geld, das als allgemeines Zahlungsmittel anerkannt ist, dient dazu, die Austauschprozesse zwischen verschiedenen Marktteilnehmern zu erleichtern: In einer „Naturalwirtschaft“ findet ein Tausch „Ware gegen Ware“ statt; der Käufer eines Gutes muss also im Zweifel dem Verkäufer genau das Gut (in geeigneter Stückelung) anbieten, das dieser gerade haben will, was in der Praxis zu unlösbaren Schwierigkeiten führen muss. In einer „Geldwirtschaft“ findet hingegen ein Tausch „Ware gegen Geld, Geld gegen Ware“ statt: Der Käufer eines Gutes gibt im Austausch Geld, das der Verkäufer dann wieder dazu verwenden kann, andere, von ihm gewünschte Waren in der präferierten Menge zu erwerben. Es ist offenkundig, dass damit Transaktions- und Suchkosten gespart werden können, arbeitsteilige Prozesse also überaus erst ermöglicht werden.

Darüber hinaus erfüllt Geld eine Funktion als Recheneinheit (1 Päckchen Butter = 1,20 Euro, aber nicht 1 Päckchen Butter = 4 Brötchen à 30 cent) und eine Funktion als Wertaufbewahrungsmittel (man kann Geld „sparen“ und damit für späteren Konsum aufheben, was mit Butter oder Brötchen nicht geht, weil diese verderben) – Geld hingegen behält seinen Wert, solange es nicht zu Inflation (=Geldentwertung) kommt.

Geld kann dabei verschiedene Formen annehmen: Bargeld (=Banknoten und Münzen) sowie Buchgeld (=Geld, das nur auf Bankkonten existiert, aber trotzdem für Zahlungszwecke eingesetzt werden kann, z.B. in Form von Überweisungen). Bargeld

kann nur durch die jeweilige Zentralbank (in Deutschland: Europäische Zentralbank bzw. vorher Deutsche Bundesbank) ausgegeben werden. Geld auf Bankkonten stellt dabei im Kern eine Forderung gegenüber der jeweiligen Bank dar, auf Wunsch Bargeld ausgezahlt zu bekommen; nur solange die Menschen „glauben“, dass dieser Forderung entsprochen wird, sind beide Formen (Buchgeld und Bargeld) äquivalent, so dass Buchgeld auch die Funktion eines Zahlungsmittels einnehmen kann. Buchgeld entsteht dabei entweder durch Überweisung eines Schuldners (z.B. unbare Lohnzahlungen des Arbeitgebers an den Arbeitnehmer) oder durch Einzahlung von Bargeld (z.B. Bildung von Sparguthaben), aber auch durch Kreditschöpfung des Bankensystems selbst (z.B. Aufnahme eines Kontokorrentkredits durch einen privaten Haushalt). Das Geldmonopol der Zentralbank ist insoweit nicht perfekt; auch Geschäftsbanken können Geld (nämlich Buchgeld) schaffen.

In manchen Regionen Deutschlands haben sich darüber regionale Geldkreisläufe mit eigenem Geld herausgebildet, die zwar typischerweise nur in der jeweiligen Region als Geld anerkannt werden, hier aber genau die gleichen Funktionen erfüllen – Ziel derartiger Initiativen ist es, regionale Wirtschaftskreisläufe anzuregen. Im Internet kursieren überdies „bitcoins“ als eine Art spezifisches Geld, das für Internetkäufe verwendet werden kann.

Die Geschäftsbanken können aber nicht unbegrenzt viel Kredit vergeben (und damit Geld schaffen), weil sie für alle auf ihren Konten liegenden Geldbestände (also Sicht- bzw. Sparguthaben ihrer Kunden bzw. Kreditvergaben – beides sind aus Sicht der Bank Verbindlichkeiten gegenüber dem privaten Sektor) eine bestimmte Menge Bargeld (das sie nicht selber emittieren können) bei der Zentralbank hinterlegen müssen (die sogenannte Mindestreserve): Bei einem Mindestreservesatz von 10% bedeutet das, dass die Geschäftsbanken für einen von ihr vergebenen Kredit von 100 Euro 10 Euro in bar bei der Zentralbank hinterlegen müssen. Die Kreditvergabekapazität der Geschäftsbanken ist daher gesamtwirtschaftlich beschränkt durch die Menge an Bargeld, das die Zentralbank ausgegeben hat (bzw. aus Sicht der einzelnen Bank: das ihr durch Einzahlungen ihrer Kunden, durch Leihe bei anderen Geschäftsbanken oder durch Leihe bei der Zentralbank zur Verfügung steht). Mit anderen Worten: Die tatsächlich umlaufende Geldmenge (bestehend aus Bargeld und geschöpften Krediten) ist zwar um ein Vielfaches höher als die Emission von Bargeld durch die Zentralbank selber; sie kann aber nicht unbegrenzt steigen.

Beispiel: Zahlt ein Kunde Bargeld in Höhe von 100 Euro bei einer Bank ein, so ergibt sich folgender Kreislauf: Die Bank muss (bei einem Mindestreservesatz von 10%) hiervon 10 Euro bei der Zentralbank hinterlegen. Es verbleiben ihr Barmittel in Höhe von 90 Euro. Sie kann also Kredite in Höhe von 900 Euro vergeben (und die vorhandenen 90 Euro an Bargeld als Mindestreserve bei der Zentralbank einzahlen). 100 Euro Bargeld, die von der Zentralbank einmal emittiert worden sind, führen in diesem Fall also zu einer Gesamtgeldmenge von 1000 Euro (100 Euro ursprüngliches Bargeld+900 Euro Buchgeld durch Kreditvergabe). Die maximal resultierende Geldmenge ergibt sich durch die Multiplikation der anfänglich dem Geschäftsbankensystem zur Verfügung gestellten Bargeldmenge (im Beispiel: 100 Euro) und dem sogenannten Geldmengenmultiplikator $m=1/mr$, mit mr =Mindestreservesatz (im Beispiel: $m=1/0,1=10$).

Auch wenn Sparguthaben (und Kredite) einen Anspruch auf die Auszahlung von Bargeld begründen, übersteigt die Summe der Kredite (und damit der potentiellen Zahlungsverpflichtung der Banken) also im Regelfall die Menge des umlaufenden Bargeldes; dies ist unbedenklich, solange nicht alle Kunden gleichzeitig ihre Bargeldforderungen einlösen wollen. Würde dies der Fall sein, wäre eine Bank unverzüglich zahlungsunfähig, weil sie Bargeld eben nicht selber emittieren kann.

Aktuelles Beispiel ist die Situation in Griechenland im Frühsommer des Jahres, als viele Bankkunden aus Sorge um eine Rückkehr zu Drachme ihre Bankguthaben auflösen wollten und in Euro-Bargeld umtauschen wollten; die Regierung reagierte hierauf mit Beschränkungen des maximalen Umtauschbetrages. In Deutschland gibt es zur Vermeidung eines solchen „bank run“ den Einlagensicherungsfonds, eine Art Beistandspflicht aller Banken, die verhindern soll, dass einzelne Banken durch derartige Vertrauensverluste ihrer Kunden in die Insolvenz getrieben werden.

Da die Zentralbank nicht die gesamte Geldmenge kontrollieren kann, sondern nur das von ihr ausgegebene Bargeld, ist ihr Einfluss auf die Geldmengenentwicklung eher indirekt: Sie kann beispielsweise die Mindestreservesätze anpassen (und damit die Kreditvergabekapazität des Geschäftsbankensystems beeinflussen), die Zinssätze verändern, zu denen sie eine Refinanzierung der Geschäftsbanken bei der Zentralbank erlaubt, oder direkt mit Käufen/Verkäufen am Wertpapiermarkt/Kreditmarkt (=Kauf/Verkauf von Wertpapieren gegen Bargeld) intervenieren. Wichtigstes geldpolitisches Instrument sind derzeit

Kreditmarktgeschäfte der Notenbank: Indem sie beispielsweise dort als Anbieter in Erscheinung tritt (also Kredite vergibt und diese mit Bargeld unterlegt), senkt sie den Zins und erhöht damit die Kreditnachfrage, so dass die gesamte Geldmenge steigt.

Man mag nun argumentieren, dass es egal ist, wie viel Geld umläuft, da für die Entscheidungsfindung von Unternehmen und Konsumenten letzten Endes nur die relativen Preise relevant sind, denn man tauscht ja Güter gegen Geld gegen Güter. Wenn für eine bestimmte Angebotsmenge (das können Güter oder auch Faktorleistungen sein) mehr Geld eingetauscht werden kann, wird man dadurch nicht „reicher“, weil man für das zu erwerbende Gut ebenfalls mehr Geld zu zahlen hat. Geld erleichtert insoweit das Wirtschaftsleben, ist für sich genommen aber neutral. Dies wird durch die sogenannte Quantitätsgleichung ausgedrückt, nach der das Ausgabenvolumen (Gütermenge Y multipliziert mit Güterpreisniveau P) der umlaufenden Geldmenge (genauer: Geldmenge M multipliziert mit der „Umlaufgeschwindigkeit des Geldes“ v) entsprechen muss

$$P \cdot Y = M \cdot v.$$

Ein Anstieg der Geldmenge würde bei gleichbleibender Produktion und konstanter Umlaufgeschwindigkeit somit lediglich das allgemeine Preisniveau erhöhen; es kommt zu Inflation.

Tatsächlich aber steigen nicht alle Preise unmittelbar in gleichem Umfang. Zum einen gibt es in vielen Bereichen vertragliche Bindungen (beispielsweise bei Löhnen oder Mieten, die nur in unregelmäßigen Abständen angepasst werden, aber auch bei längerfristigen Liefervereinbarungen), so dass sich die Mehrnachfrage aufgrund zusätzlichen Geldes zunächst auf die übrigen Märkte konzentriert (und hier unter Umständen überproportionale Preissteigerungen auslöst, die wiederum zu einer Angebotsausweitung führen). Zum anderen verteilt die Zentralbank das zusätzliche Geld ja nicht an alle Wirtschaftssubjekte in gleicher Weise, sondern interveniert selektiv über ihre geldpolitischen Instrumente (in erster Linie: Kreditmarktgeschäfte): Es profitieren also in diesem Fall zunächst einmal die unmittelbaren Kreditnehmer, die nun über eine höhere Kaufkraft verfügen. Es kommt daher zu einer Mehrnachfrage zunächst bei denjenigen Gütern, die von diesen Kreditnehmern bevorzugt nachgefragt werden (dies sind im Zweifel am ehesten Unternehmen, die kreditfinanzierte Investitionen tätigen wollen). Die Mehrnachfrage bewirkt dann auf diesen Märkten im

Regelfall auch steigende Produktion und steigende Preise, wirkt also „expansiv“. Das aber bedeutet, dass auch Geldpolitik durchaus realwirtschaftliche Auswirkungen haben kann, also nicht notwendigerweise neutral ist. Schließlich ist es denkbar, dass es auf allen relevanten Gütermärkten unterausgelastete Kapazitäten gibt, die Nachfrage also kleiner ist als das (potentielle) Angebot: Kommt dann mehr Geld in Umlauf, steigt die Nachfrage und die unterausgelasteten Kapazitäten werden besser ausgelastet, ohne dass es bereits zu Preissteigerungen kommt. Über diese Mechanismen kann die Geldpolitik daher auch realwirtschaft intervenieren.

Ein Anstieg des allgemeinen Preisniveaus (=Inflation) kann zwar durch Änderung einzelner Preise verursacht sein (z.B.: Energiepreise, die als Kostenfaktor in viele weitere Preise eingehen); auch allgemeine Lohnsteigerungen können inflationär wirken. Letzten Endes setzt eine Inflation aber immer eine entsprechende Ausweitung der Geldmenge voraus (umgekehrt gilt das, wie gezeigt, aber nicht), weshalb die meisten Zentralbanken ihre hauptsächliche Aufgabe darin sehen, die Preisstabilität zu bewahren (und deswegen die Geldversorgung „knapp“ zu halten).

Eine „moderate“ Inflation ist nach allgemeiner Auffassung unschädlich, ja sogar notwendig, weil Preise „nach unten“ zumeist starr sind, Änderungen der relativen Preise (die im Zuge allgemeiner Marktprozesse immer notwendig sein werden) also nur durch unterschiedlich starke Preissteigerungen erreicht werden können. Die Europäische Zentralbank strebt daher eine Inflationsrate von 2% pro Jahr an und ist bemüht, die Geldmenge in einer Weise auszuweiten, dass genau dieser Wert erreicht wird (indem die Geldmenge im Umfang von gewünschter Inflationsrate=2% und erwarteter Veränderung der realen Produktion=Wirtschaftswachstum erhöht wird). Eine hohe (oder gar sich beschleunigende=„galoppierende“) Inflation ist allerdings mit erheblichen Verwerfungen und Verelendung breiter Bevölkerungsschichten verbunden – schlechtes Beispiel aus der deutschen Vergangenheit ist die Hyperinflation des Jahres 1923, in deren Endphase sich die Preise innerhalb von einer Woche verzehnfachten, was schließlich eine Währungsreform notwendig machte). Eine Deflation, also ein sinkendes Preisniveau, ist allerdings auch unerwünscht, weil sinkende Preise durch sinkende Nachfrage ausgelöst werden; analog zu den oben beschriebenen Mechanismen kann das dann negative realwirtschaftliche Folgen haben (Unterbeschäftigung; soziale Verwerfungen).

Die Geldpolitik kann überdies auch eingesetzt werden, um den „Außenwert“ einer Währung zu beeinflussen, also den Wechselkurs. Bis in die 1970er Jahre hinein galten weltweit feste Wechselkurse; so war der Kurs des US-Dollar auf 4 DM festgelegt. Die Notenbanken waren verpflichtet, diesen Kurs zu stabilisieren: Bei einer Mehrnachfrage nach Dollar (z.B. aufgrund einer steigenden Nachfrage nach US-amerikanischen Produkten), die für sich genommen den Dollarkurs (bzw. –kurs) erhöht hätte, war die amerikanische Notenbank (die Federal Reserve Bank) verpflichtet, DM zum festgelegten Kurs anzukaufen und damit die nachgefragte Dollarmenge zur Verfügung zu stellen; Folge wäre in diesem Fall eine Ausweitung der amerikanischen Geldmenge und gleichzeitig eine Verringerung der umlaufenden DM-Menge. Umgekehrt wäre auch die Deutsche Bundesbank verpflichtet gewesen, Dollar zum festgelegten Kurs zur Verfügung zu stellen, was aber nur möglich ist, wenn die Bundesbank über entsprechende Devisenreserven verfügt.

Dieses System ist im globalem Maßstab allerdings Anfang der 1970er Jahre zusammengebrochen (weil die Notenbanken in diesem System die Kontrolle über die heimische Geldmenge verloren haben) und durch ein System flexibler Wechselkurse ersetzt. In diesem Fall würde eine Mehrnachfrage nach US-Dollar den Dollarkurs erhöhen (aber die Geldmenge in beiden beteiligten Ländern unverändert lassen). Flexible Wechselkurse erlauben daher in stärkerem Maße als feste Wechselkurse eine eigenständige Geldpolitik in den beteiligten Währungsräumen und sind insoweit aus nationaler Sicht zu bevorzugen.

In den 1970er Jahren wurden dann innerhalb Europas erneut feste Wechselkurse eingeführt, um auf diese Weise die europäische Integration zu stärken (feste Wechselkurse erhöhen die Planungssicherheit für grenzüberschreitenden Handel, weil keine wechselkursbedingten Preisänderungen zu erwarten sind). Dieses „Europäische Währungssystem“ kann als der Vorläufer des Euro gelten, der Ende der 1990er Jahre in den meisten Mitgliedsländern der Europäischen Union eingeführt wurde.

Gegenüber den meisten anderen Währungen (vor allem: US-Dollar und britisches Pfund) gelten weiterhin flexible Wechselkurse; allerdings greifen die Notenbanken häufig mit Devisenkäufen/-verkäufen in die Preisbildung ein, um allzu starke Schwankungen der Wechselkurse zu vermeiden.

3. Grundzüge der keynesianischen Makroökonomik

a. Grundlagen

Im Folgenden sollen jetzt die Grundzüge der (keynesianischen) Makroökonomik dargestellt werden – auf Formeln kann dabei nicht verzichtet werden; diese sind aber bewußt einfach gehalten.

Die Makroökonomik unterscheidet nach Sektoren (Haushalte, Unternehmen, Staat, Ausland=Rest der Welt) und nach Nachfrageaggregaten (Konsum, Investitionen, Staatsnachfrage, Auslandsnachfrage=Export). Jedem Sektor ist ein Nachfrageaggregat zugeordnet: Haushalte sind Träger der Konsumnachfrage, Unternehmen sind Träger der Investitionsnachfrage, der Staat fragt Güter zur Bereitstellung öffentlicher Leistungen („Staatsverbrauch“) nach, das Ausland fragt definitionsgemäß nur Exporte nach. Auslandsnachfrage und Staatsnachfrage können sich natürlich auf Investitionsgüter oder Konsumgüter richten, aber diese Unterscheidung wird normalerweise auf makroökonomischer Ebene nicht vorgenommen.

Das aggregierte Angebot wird von den Unternehmen des Inlands bzw. durch das Ausland (=Importe) bereitgestellt. Die inländischen Unternehmen greifen dabei auf verschiedene Produktionsfaktoren (vereinfacht: Arbeit L und Kapital K) zurück; es gibt also wieder eine Produktionsfunktion, die gleiche Eigenschaften hat wie im mikroökonomischen Teil dargestellt. Als Formel ausgedrückt bedeutet dies:

$$Y^s = f(L, K),$$

mit Y^s =inländisches Angebot (mit dem Index s für „supply“), L=Arbeitseinsatz, K=Kapitaleinsatz, $f(*)$ =Produktionsfunktion. In gewisser Weise wäre dies also eine „entstehungsseitige“ Betrachtung der Erstellung des Bruttoinlandsprodukts Y.

Beide Produktionsfaktoren werden wiederum von den privaten Haushalten bereitgestellt (bei Arbeit ist das offensichtlich, bei Kapital ist die Annahme, dass Unternehmen ihren Kapitaleinsatz durch Aufnahme von Krediten finanzieren, hierfür also die Ersparnis der privaten Haushalte in Anspruch nehmen).

Aus didaktischen Gründen sei zunächst vereinfacht eine geschlossene Volkswirtschaft betrachtet: Es gibt also kein Ausland. Diese vereinfachende Annahme dient lediglich der Veranschaulichung grundlegender Zusammenhänge und wird später wieder

aufgehoben. Aus den oben gemachten Annahmen folgt dann: Das im Produktionsprozess entstehende Einkommen Y fließt allein den privaten Haushalten zu, sei es als Arbeitseinkommen $w \cdot L$ oder als Gewinn- bzw. Vermögenseinkommen $r \cdot K$. Es gilt also

$$(1) \quad Y = w \cdot L + r \cdot K$$

mit w =Lohnsatz („wage“), r =Zins („rate of return“). Man kann (1) auch als „verteilungsseitige“ Betrachtung verstehen, den es wird gezeigt, wie sich das gesamtwirtschaftliche Einkommen auf die beiden Produktionsfaktoren aufteilt. Zu beachten ist dabei, dass der Sektor der Unternehmen in institutioneller Betrachtung kein eigenes Einkommen erzielt (sondern nur die Unternehmenseigner bzw. die Kapitalgeber).

Die von den privaten Haushalten erzielten Einkommen (Y) können nun in mehrfacher Weise verwendet werden: Sie fließen entweder an den Staat (als Steuerzahlungen T), werden von den privaten Haushalten selber konsumiert (Konsum C), oder sie werden gespart (Ersparnis S). Es gilt also:

$$(2) \quad Y = C + T + S.$$

Dies ist eine reine Definitionsgleichung und ist unabhängig von irgendwelchen Verhaltensannahmen. Da das Einkommen Y nur im Produktionsprozess entsteht, ist das Y in Gleichung (2) zugleich immer gleich groß wie das gesamtwirtschaftliche Angebot Y^s aus der Formel oben.

Die gesamtwirtschaftliche aggregierte Nachfrage Y^d (mit dem Index d für „demand“) ist wiederum verwendungsseitig ganz offenkundig definiert als

$$(3) \quad Y^d = C + I + G.$$

Während Gleichung (1) die Verteilung des Volkseinkommens widerspiegelt, gibt Gleichung (3) dessen Verwendung an.

Es gilt jetzt, dass erzielte Einkommen Y , aggregiertes Güterangebot Y^s und aggregierte Nachfrage Y^d sich stets einander entsprechen müssen. Es ergibt sich also die folgende „Identität“

$$(4) \quad Y = Y^s = Y^d \Leftrightarrow C + T + S = C + I + G$$

bzw.

$$(5) \quad I + G = T + S.$$

Die Summe aus Investitionen I und Staatsnachfrage G entspricht also immer der Summe aus privater Ersparnis S und Steueraufkommen T . Wenn sich der Staat nicht verschulden darf (die Staatsnachfrage also nur aus Steuern finanziert werden kann), reduziert sich dies sogar zu der Gleichung

$$(6) \quad I = S.$$

Gleichung (4) ist die grundlegende makroökonomische Beziehungsgleichung. Während (1) bis (3) noch reine Definitionsgleichung waren, postulieren (4) bis (6) Gleichgewichtsbedingungen. Diese sind zwar ex post immer erfüllt, nicht notwendigerweise aber ex ante: Denkbar sind Fälle, in denen die (geplante) Nachfrage geringer (höher) ausfällt als das (geplante) Angebot, was nach Gleichung (5) bzw. (6) gleichbedeutend ist mit einer positiven (negativen) Diskrepanz zwischen (geplanter) Ersparnis und (geplanten) Investitionen. Ex post ergibt sich in diesem Fall eine Produktionseinschränkung (bzw. eine ungeplante Mehrersparnis: Die kürzere Marktseite setzt sich immer durch), so dass die Gleichung (4) bis (6) weiterhin erfüllt sind; dieses Gleichgewicht wird aber ex ante nicht notwendigerweise automatisch erreicht.

Im Kern handelt es sich hierbei um die Kontroverse zwischen (klassischer bzw. neoklassischer) „Angebotstheorie“ und keynesianischer „Nachfragetheorie“. Die Denkschule der Neoklassik geht davon aus sich auch auf aggregierter Ebene immer ein Gleichgewicht einstellt, weil es zu Zinsanpassungen kommt: Ein Mehrangebot an Gütern wäre gleichbedeutend mit einem höheren Einkommen und damit mit einem Überangebot an Sparkapital; in der Folge würden die Zinsen sinken und damit die Investitionsgüternachfrage solange anregen, bis das anfängliche Überangebot an Gütern tatsächlich nachgefragt wird (im Falle eines zu geringen Güterangebots umgekehrt). („Say'sches Gesetz“, nach Jean Baptiste Say (1767-1832): „Jedes Angebot schafft sich seine Nachfrage“). Die keynesianische Theorie geht demgegenüber davon aus, dass in bestimmten Situationen dieses neuerliche Gleichgewicht nicht erreicht wird und deswegen staatliche nachfragestützende Maßnahmen erforderlich sein können. Dies ist beispielsweise dann der Fall, wenn die Zinsen bereits sehr niedrig sind oder die Unternehmen so ungünstige

Gewinnerwartungen haben, dass sie auch durch niedrige Refinanzierungskosten nicht zu vermehrten Investitionen angeregt werden.

Aus Gleichung (4) lässt sich zudem ableiten, dass „autonome“ Nachfrageimpulse Effekte haben können, die deutlich über den anfänglichen Impuls hinausgehen. Definiert man nämlich

$$(7) \quad c = C/Y \text{ (Konsumquote)}$$

$$t = T/Y \text{ (Steuerquote)}$$

$$s = S/Y \text{ (Sparquote)}$$

so lässt sich (2) auch umformulieren zu

$$(8) \quad Y = c Y + t Y + s Y$$

$$\Rightarrow C = c Y = (1-s-t) Y$$

Einsetzen in (3) führt zu

$$(9) \quad Y = (1-s-t) \cdot Y + I + G$$

$$\Rightarrow Y - (1-s-t) \cdot Y = I + G$$

$$\Rightarrow Y \cdot (1-(1-s-t)) = I + G$$

$$\Rightarrow Y = 1/(s+t) (I + G)$$

Güterangebot und –nachfrage übersteigen also die Summe aus Investitionen und Staatsnachfrage um ein Vielfaches, denn der Bruch in Gleichung (9) ist stets größer als 1: Bei einer Sparquote von 10% und einer Steuerquote von 20% ist nämlich $1/(s+t)=1/0,3=3,33$. Der Bruch $1/(s+t)$ wird deshalb auch Einkommensmultiplikator genannt.

Hieraus folgt, dass auch kleine Veränderungen der beiden genannten Nachfrageaggregate hohe Einkommenseffekte auslösen können. Es gilt nämlich

$$(10) \quad \Delta Y = 1/(s+t) (\Delta I + \Delta G)$$

mit Δ =Veränderung der jeweils nachfolgend genannten Größe.

Zahlenbeispiel

Angenommen seien (wie oben) Werte für die Spar- und die Steuerquote von 0,1 bzw. 0,2. Betrachtet wird nun eine Erhöhung der Staatsnachfrage G um 1.000 Euro.

Im ersten Schritt erhöht sich dadurch die Gesamtnachfrage um eben jene 1.000 Euro. Hiervon werden 100 Euro gespart und 200 Euro als zusätzliche Steuereinnahmen an den Staat abgeführt; die übrigen 700 Euro werden konsumiert und erhöhen das Einkommen. Von dem zusätzlichen Einkommen werden wiederum 30% zu Ersparnis bzw. zu Steuereinnahmen, 70% werden erneut konsumiert (und erhöhen das Einkommen um $0,7 \cdot 700 \text{ Euro} = 490 \text{ Euro}$. Hiervon werden 70% konsumiert, das Einkommen steigt um weitere $0,7 \cdot 490 = 343 \text{ Euro}$ – dieser Prozess setzt sich fort, bis zum Schluss ein kumulierter Einkommensanstieg um $1/0,3 \cdot 1000 = 3.333 \text{ Euro}$ erreicht ist.

Annahme dabei ist: Die zusätzliche Ersparnisbildung „versickert“, führt also nicht zu zusätzlichen Investitionen; gleiches gilt auch für die zusätzlichen Steuereinnahmen. Man mag diese Annahme für unrealistisch halten; der Einkommensmultiplikator soll aber vor allem aufzeigen, wie in einer Situation dauerhafter Nachfrageschwäche durch „autonome“ Mehrnachfrage die Wirtschaft angeregt werden kann.

Die Einkommenswirkungen treten natürlich nur ein, wenn die Finanzierung der zusätzlichen Nachfrage (I bzw. G) nicht durch zusätzliche Ersparnisbildung oder durch Steuern erfolgt, sondern wenn diese mit dem Zufluss zusätzlicher Mittel verbunden ist – mit anderen Worten, wenn es sich um eine zusätzliche Verschuldung handelt, die nicht zu Veränderungen der Steuer- bzw. Sparquote führt.

Im Falle einer offenen Volkswirtschaft bleiben die genannten Zusammenhänge weitgehend erhalten. Die Verwendungsgleichung (3) muss in diesem Fall ergänzt werden um den sogenannten Außenbeitrag X-M (X=Export, M=Import):

$$(11) \quad Y^d = C + I + G + X - M$$

Ein Teil der Nachfrage richtet sich auf Importgüter (M) und wird daher nicht im Inland nachfragewirksam. Auf der anderen Seite erhöht sich aber die Gesamtnachfrage um die Nachfrage nach Auslands nach inländischen Güter (X). Die nachfolgenden Gleichungen (4) bis (9) müssen jetzt ebenfalls um den Außenbeitrag erweitert werden; für den Einkommensmultiplikator ergibt sich dann:

$$(12) \quad \Delta Y = 1/(s+t+m) (\Delta I + \Delta G + \Delta X)$$

mit $m=M/Y$ =Importquote

Der Einkommensmultiplikator ist im Fall der offenen Volkswirtschaft kleiner als im Fall der geschlossenen Volkswirtschaft, weil ein Teil der zusätzlichen „autonomen“ Nachfrage via Importe ins Ausland abfließt. Bei den autonomen Nachfragekomponenten (zweiter Klammerausdruck in Gleichung (12)) muss aber natürlich auch der zusätzliche Export mit aufgenommen werden; Multiplikatorprozesse können also nicht nur durch heimische Investitionen bzw. heimische Staatsnachfrage, sondern auch durch das Ausland ausgelöst werden – mit ein Grund dafür, dass in aktuellen politischen Diskussionen häufig gefordert wird, Deutschland solle mehr importieren, um auf diese Weise die wirtschaftliche Entwicklung im Rest Europas anzuregen.

b. Konjunkturpolitische Schlussfolgerungen

Aus den voranstehenden Ausführungen folgt, dass im Falle unerwünschter konjunktureller Schwankungen vor allem die Finanzpolitik eine aktive Rolle einnehmen sollte: Bleibt die gesamtwirtschaftliche Nachfrage hinter den Angebotsmöglichkeiten zurück, kommt es zu Unterbeschäftigung, die bei unzureichender Flexibilität von Preisen, Zinsen und Löhnen durch Marktanpassungen nicht abgebaut wird – im Gegenteil, weil Unterbeschäftigung mit niedrigeren Einkommen verbunden ist (Arbeitslose erhalten ja bestenfalls eine Arbeitslosenunterstützung), kann es sogar zu dauerhafter Unterauslastung der Kapazitäten kommen. In einer solchen Situation kann der Staat durch defizitfinanzierte Mehrnachfrage die anfängliche Nachfrangelücke schließen und infolge des Multiplikatoreffekts sogar mit verhältnismäßig geringem Mitteleinsatz hohe Impulse erzielen. Dies ist der Kern keynesianischer Nachfragepolitik zur Konjunktursteuerung: In einer Rezession (=Unterauslastung der Kapazitäten) solle der Staat nachfragestützend eingreifen und so die Konjunktur stabilisieren („Fiskalpolitik“); die Finanzierung muss dabei durch staatliche Kreditaufnahme erfolgen.

Auch die Geldpolitik kann zur Konjunkturstabilisierung eingesetzt werden, indem in einer konjunkturellen Schwächephase die Geldmenge ausgeweitet wird (bzw. die Notenbankzinsen gesenkt werden); allerdings sind die Wirkungen möglicherweise nur schwach, weil der Mechanismus eher indirekt ist: Zinssenkungen sollen die

Investitionen anregen; wenn aber die Gewinnerwartungen der Unternehmen ungünstig sind, werden auch niedrigere Zinsen nicht den erwünschten expansiven Effekt haben. Hinzu kommt, dass die Möglichkeiten einer Zinssenkung beschränkt sind, weil der Zins unter normalen Umständen nicht unter Null fallen kann. Deswegen wird eher die Fiskalpolitik als Mittel der Wahl angesehen, um aktive Konjunkturpolitik zu betreiben.

Gesetzlich geregelt ist dies im „Gesetz zur Förderung der Stabilität und des Wachstums der Wirtschaft“ („Stabilitäts- und Wachstumsgesetz“) von 1967:

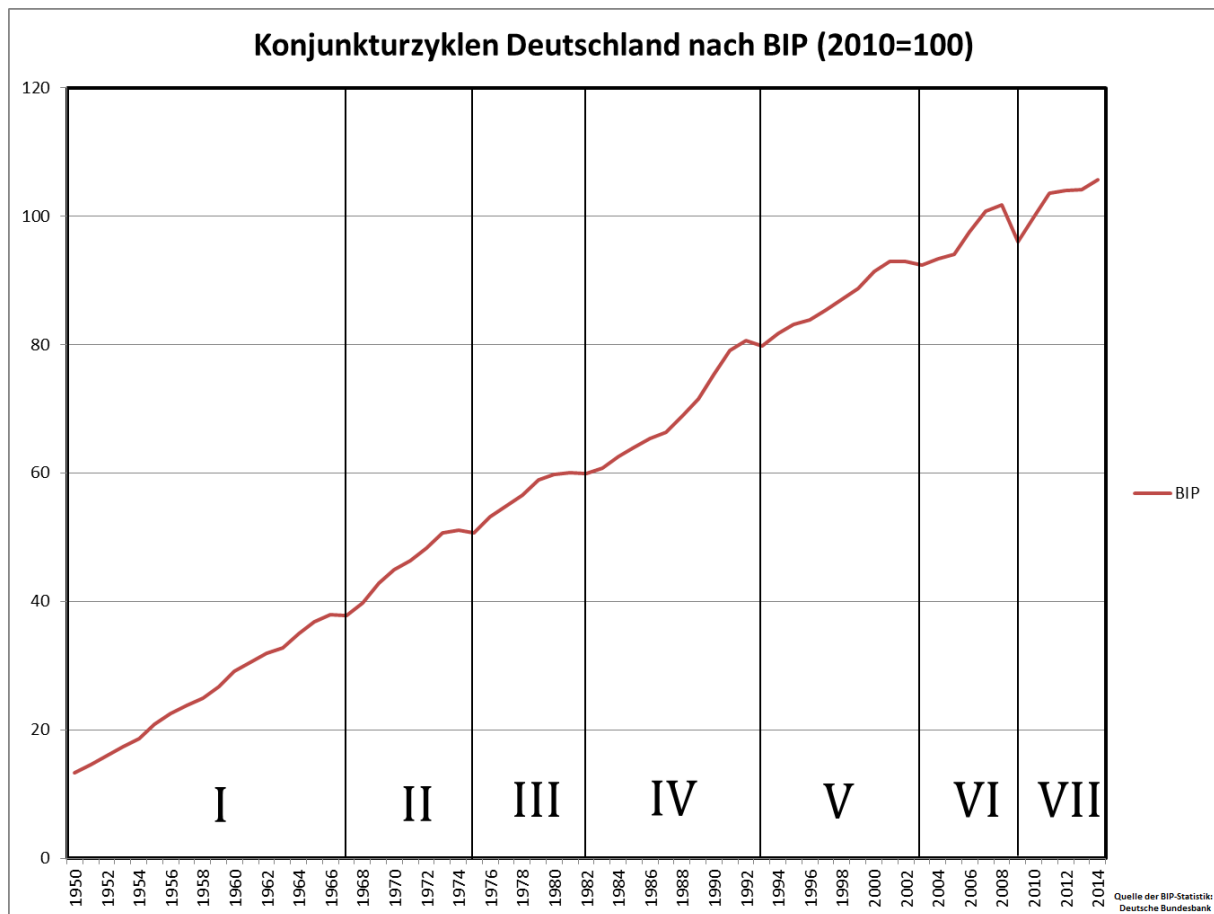
§1: Bund und Länder haben bei ihren wirtschafts- und finanzpolitischen Maßnahmen die Erfordernisse des gesamtwirtschaftlichen Gleichgewichts zu beachten. Die Maßnahmen sind so zu treffen, daß sie im Rahmen der marktwirtschaftlichen Ordnung gleichzeitig zur Stabilität des Preisniveaus, zu einem hohen Beschäftigungsstand und außenwirtschaftlichem Gleichgewicht bei stetigem und angemessenem Wirtschaftswachstum beitragen.

Das Gesetz definiert das sogenannte „magische Viereck“ der Wirtschaftspolitik (Preisstabilität, hohe Beschäftigung, außenwirtschaftliches Gleichgewicht, stetiges und angemessenes Wirtschaftswachstum), wobei der Auftrag der Konjunkturstabilisierung durch den Begriff „stetiges Wirtschaftswachstum“ abgebildet wird. Das Stabilitätsgesetz konkretisiert in seinem weiteren Text dann die Maßnahmen, mit denen diese Ziele erreicht werden sollen; unter anderem durch eine konjunkturorientierte Ausgaben- und Einnahmenpolitik.

Die übrigen Ziele sind weitgehend selbsterklärend; das Ziel des „außenwirtschaftlichen Gleichgewichts“ ist nur aus der Zeit heraus verständlich (1967: feste Wechselkurse, die gegen Veränderungen abgesichert werden sollten) und ist heute nicht weiter bedeutsam (Wechselkursstabilisierung, das Pendant des „außenwirtschaftlichen Gleichgewichts“, wird nicht als Aufgabe der Politik angesehen). Der Begriff „magisches Viereck“ wurde geprägt, weil diese Ziele teilweise in Konflikt miteinander stehen können.

Das folgende Schaubild zeigt die Entwicklung des Bruttoinlandsprodukts (Y) in Deutschland von 1950-2014. Erkennbar ist, dass es zwar mehrere Phasen schwacher wirtschaftlicher Aktivität gab (1967/68, 1974/75, 1980-1982, 1992/93, 2001/02 und 2008/09), dass aber wirklich massive Rezessionen im Sinne eines länger andauernden Rückgangs des Bruttoinlandsprodukts die Ausnahme waren. Vielmehr

waren die meisten konjunkturellen Schwächephasen eher durch nachlassende Wachstumsraten des Bruttoinlandsprodukt, jedoch nicht durch einen absoluten Rückgang gekennzeichnet. Häufig wird daher nicht die Entwicklung des Bruttoinlandsprodukts (bzw. die Wachstumsrate des Bruttoinlandsprodukts), sondern vielmehr der Auslastungsgrad der vorhandenen Produktionskapazitäten (ermittelt als Differenz zwischen angebotsseitigem Produktionspotential und Gesamtnachfrage) herangezogen.



Gegenstück einer konjunkturellen Schwäche ist eine konjunkturelle Überhitzung, bei der die Nachfrage stark steigt (bzw. die Kapazitäten voll ausgelastet sind). Problematisch ist dies deswegen, weil es in einer solchen Situation zu überproportionalen Preissteigerungen kommen kann, also zu inflationären Tendenzen. Derartige Situationen gab es beispielsweise in den Jahren 1971/72 (ausgelöst durch überproportionale Lohnsteigerungen) oder in den Jahren 1990/91 (ausgelöst durch den Nachfrageboom aus Ostdeutschland); in diesem Fall wären Geld- und Fiskalpolitik in Analogie zu expansiven Maßnahmen im Fall einer konjunkturellen Schwächephase gefordert, durch Zinserhöhungen bzw. durch temporäre

Ausgabekürzungen/Steuererhöhungen die gesamtwirtschaftliche Nachfrage zu dämpfen. Das Stabilitäts- und Wachstumsgesetz sieht auch hierfür eine entsprechende Verpflichtung vor. Allerdings zeigt die Erfahrung, dass konjunkturdämpfende Maßnahmen von der Politik sehr viel halbherziger realisiert werden als konjunkturstimulierende Maßnahmen.

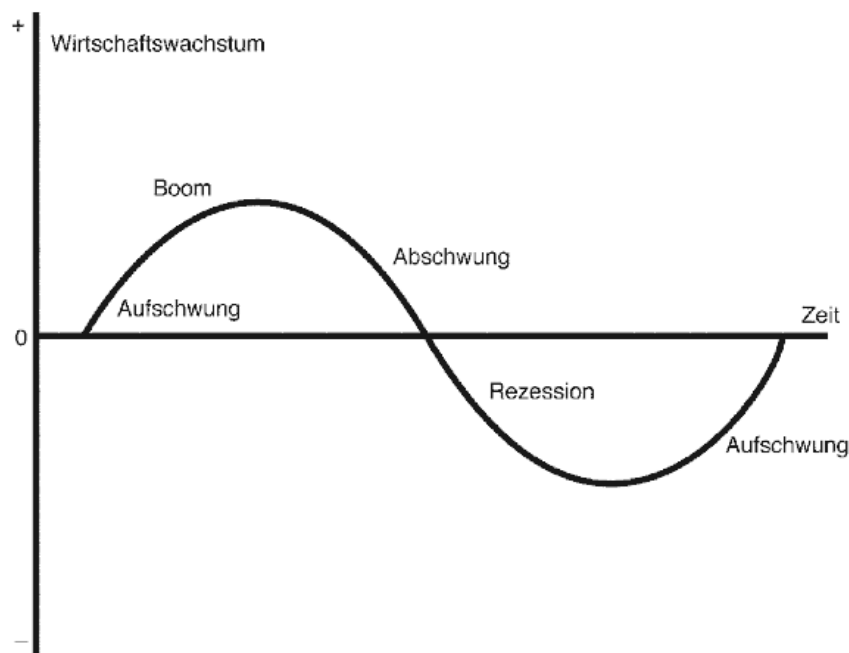
Fiskalpolitische Maßnahmen zur Stabilisierung der Konjunktur kamen vor allem in den Jahren 1967/68 und 1974/75 sowie aktuell wieder 2008/2009 zur Anwendung; zwischenzeitlich hatte sich die Politik von keynesianischen Vorstellungen zur Stabilisierung der Wirtschaft abgewandt. Grund dafür ist, dass eine keynesianische Nachfragepolitik eine Reihe von praktischen Problemen aufweist (zeitliche Steuerung; Asymmetrie von expansiven und kontraktiven fiskalpolitischen Maßnahmen mit der Folge steigender Staatsverschuldung).

Es stellt sich hier die Frage, wie es überhaupt zu derartigen Konjunkturschwankungen kommen kann. Die makroökonomische Theorie unterscheidet zwischen sehr verschiedenen Zyklen, die eine Volkswirtschaft durchlaufen kann:

- Langfristige Zyklen („Kondratieff-Zyklen“) umfassen sehr lange Zeiträume, typischerweise 50 Jahre und mehr. Sie werden ausgelöst durch grundlegende technologische Neuerungen, die nach und nach alle Wirtschaftsbereiche durchdringen und auf diese Weise zu starken Produktivitätssteigerungen (und dadurch ausgelöste Lohn-, Einkommens- und Nachfragesteigerungen) führen. Beispiele für derartige grundlegende Innovationen sind die Erfindung der Dampfmaschine und ihre Verwendung vor allem in der Textilindustrie (ab 1780); die Erfindung und Verbreitung der (dampfmaschinengetriebenen) Eisenbahn (ab 1840); die Erfindung und Verbreitung von elektro- und benzingetriebenen Motoren sowie Innovationen in der Chemieindustrie (z.B. Kunstdünger) (ab 1880), die Erfindung und Verbreitung von elektronischen Bauelementen (ab etwa 1940) und – möglicherweise – die Weiterentwicklungen in der Informations- und Kommunikationstechnologie (ab 1980).

Davon zu unterscheiden sind Konjunkturschwankungen (mit einem Zyklus von ungefähr 5-10 Jahren). Diese weisen typischerweise das in der nachstehenden Abbildung gezeigte Verlaufsmuster des Wirtschaftswachstums auf: Ausgehend von einem konjunkturellen Aufschwung (I) eine Boomphase (II), eine Abschwungphase (III), eine Rezession (IV). Nicht jeder Konjunkturzyklus ist

aber komplett in dem Sinne, dass alle vier Phasen in gleicher Weise auftreten. Analoge konjunkturelle Zyklen lassen sich auch bei anderen wirtschaftlichen Indikatoren (insbesondere Preissteigerungsrate und Arbeitslosigkeit finden).



Konjunkturschwankungen der beschriebenen Art werden häufig durch externe Schocks ausgelöst werden (z.B. Naturkatastrophen, Börsenturbulenzen, Rohstoffpreissteigerungen, Entwicklungen im Ausland o.ä.). Diese können sich jedoch verstärken, wenn sie eine hinreichend große Intensität („Schockwirkung“) aufweisen oder zu einer verstärkten Verunsicherung der privaten Wirtschaftsakteure beitragen. Häufig werden Konjunkturschwankungen auch durch die Wirtschaftspolitik selber ausgelöst (oder verstärkt), so weil das „timing“ konjunkturpolitischer Maßnahmen wegen verschiedener Erkenntnis- und Wirkungsverzögerungen, aber auch wegen politisch gewünschter „Wahlgeschenke“ o.ä. nicht konjunkturneutral ist. Konjunkturschwankungen können aber auch „endogen“ ausgelöst werden, sind insoweit eine nahezu unvermeidliche Konsequenz wirtschaftlicher Entwicklung.

Zum Beispiel führen gewinngetriebene Investitionen zunächst zwar zu einem Nachfrageboom; da Investitionen aber immer auch einen Kapazitätseffekt aufweisen, also zusätzliches Produktionspotential schaffen, kommt es mit hoher Wahrscheinlichkeit zu einer Situation, in der die zusätzlichen Kapazitäten nicht ausgelastet werden, also die Notwendigkeit zu weiteren Investitionen

schwindet, also die Nachfrage rückläufig ist, also die wirtschaftliche Dynamik sich insgesamt abschwächt – diese Situation würde ohne konjunkturpolitische Eingriffe erst dann beendet, wenn aufgrund von Ersatzbedarfen bei „schrottreifen“ Investitionsgütern eine neuer Investitionsaufschwung in Gang käme.

- Schließlich gibt es noch sehr kurzfristige, typischerweise aber regelmäßige Schwankungen der wirtschaftlichen Aktivität, die durch Saisoneffekte oder Kalenderunregelmäßigkeiten ausgelöst werden. So kann in der Bauwirtschaft oder der Landwirtschaft häufig im Winter witterungsbedingt nicht gearbeitet werden, so dass in diesen Branchen in den Wintermonaten die Produktion deutlich geringer ausfällt als in den Sommermonaten (in denen die Produktionsausfälle im Winter dann nachgeholt werden). Ähnliches findet sich auch im Handel (Umsatzsteigerungen im Vorweihnachtsgeschäft) oder im Gastgewerbe (Umsatzsteigerungen zur Ferienzeit im Sommer). Derartige Schwankungen sind für die mittelfristige wirtschaftliche Entwicklung ohne größere Bedeutung.