



The epistemic benefits of generalisation in modelling I: Systems and applicability

Aki Lehtinen¹

Received: 5 December 2020 / Accepted: 4 June 2021 / Published online: 25 June 2021
© The Author(s), under exclusive licence to Springer Nature B.V. 2021

Abstract

This paper provides a conceptual framework that allows for distinguishing between different kinds of generalisation and applicability. It is argued that generalising models may bring epistemic benefits. They do so if they show that restrictive and unrealistic assumptions do not threaten the credibility of results derived from models. There are two different notions of applicability, generic and specific, which give rise to three different kinds of generalizations. Only generalising a result brings epistemic benefits concerning the truth of model components or results. Abstracting the model and applying the model into new systems are not intrinsically epistemically beneficial in this way. The Dixit-Stiglitz model of monopolistic competition is used as an illustration.

Keywords Generalisation · Epistemology of modelling · Model applicability · Systems · Dixit-Stiglitz model

1 Introduction

Generalising models is considered important in mathematical modelling. Modellers routinely use several expressions to indicate an epistemic benefit from generalisation. In economics, for example, they may say that an assumption has been *relaxed* or that they have *dropped restrictive assumptions*. Philosophers of economics have long recognised such epistemic benefits. Bert Hamminga (1983) observed, for example, that ‘it is widely accepted by economists that the less restrictive assumptions are needed, the higher the “value” of the result’.

However, there is no philosophical account of generalisation that explains precisely what kind of generalisation provides epistemic benefits. A generalisation yields an *epistemic benefit* if it justifiably increases the modellers’ confidence in the truth of a model component or result. Some generalisations do not provide epistemic

✉ Aki Lehtinen
aki.lehtinen@helsinki.fi

¹ College of Philosophy, Nankai University, Tianjin, China

benefits, however, and this is why some philosophers (e.g., Potochnik, 2017) have understandably found it difficult to justify generality via its contribution to truth. The purpose of this paper is to provide an account that is capable of distinguishing between different kinds of generalisations, and of singling out the cases in which it is indeed truth-conducive. Only epistemic benefits that pertain to *modifications* to a model that already applies to some system are considered, and I do not aim to provide a comprehensive account of the epistemic benefits of models.

Increasing a model's generality in an epistemically beneficial way helps to solve some of the problems that arise from the necessity of making unrealistic assumptions. Tractability considerations tend to imply that modellers describe their targets by means of assumptions that are less general than they think can really be asserted about them. In many cases, therefore, they are able to prove a result only for a special case. They know that specificities sometimes introduce falsehoods, yet they do not know whether the results of the model crucially depend on them. When the model is generalised but continues to imply the same result, the particular falsehoods are shown not to be responsible for the result. Obtaining the same result with less restrictive assumptions increases the modellers' confidence that it is not an artefact of specific assumptions known to be unrealistic. The epistemic importance of generalisations thus derives from the same kind of considerations that Kuorikoski et al. (2010) claim to motivate derivational robustness analysis (see also R  z, 2017), and herein lies its main epistemic advantage: They guard against error in showing that the results do not depend on particular falsehoods.

Generality is typically defined as 'the property of applying widely', and is measured either by the number of phenomena a model can explain or predict, or by the number of systems or targets to which it applies (Levins, 1966, 1993; Lewis & Belanger, 2015; Matthewson & Weisberg, 2009; Weisberg, 2004, 2013). Such definitions have been developed in the recent philosophical discussion on generalisation, which nevertheless has not been framed in terms of epistemic benefits.

The notion of generalising a model depends on what one means by the applicability and the application of a model. I will show that there are two different notions of model applicability, *generic* and *specific*. Many systems share properties with other systems, and 'generalised modelling' (see Weisberg, 2013, pp. 114–121) aims at describing such shared properties. A model applies *generically* to a system if the modeller aims to account for such properties, and the model is able to describe them. A model applies to a system *specifically* if it is able to describe some properties specific to it. The conditions for successful applicability are quite different in the two cases.

We will see that the epistemic benefits pertain to increasing generic rather than specific applicability. Demonstrating greater specific applicability may have other benefits that some scholars might want to call 'epistemic' in a broader sense. For example, they might be taken to increase explanatory power (e.g., Strevens, 2008; Woodward & Hitchcock, 2003), unifying power (e.g., Kitcher, 1989), cognitive economy, or testability of a theory. This paper will not discuss the relative importance of such other benefits compared to the strictly truth-related epistemic benefits, nor will there be further discussion on the concept of epistemic benefit. The aim is rather to provide a conceptual framework which helps identifying strictly truth-related epistemic benefits from generalisation.

The account developed in this paper can be used to elucidate philosophical discussions beyond the epistemic benefits in modelling, however, because discussing other benefits from generalisation requires recognising that there are three different kinds of them, and not every benefit is associated with every kind of generalisation. For example, different accounts of explanation rely on different kinds of generalisation because of their connection to the depth of explanations and their unifying power. Judgments about the trade-offs in modelling (e.g., Matthewson & Weisberg, 2009; Orzack & Sober, 1993) also depend on the notion of generalisation. However, developing these suggestions further is left for future work.

The two notions of applicability enable distinguishing between three different kinds of generalisation: (a) *applying a model to a new system*, (b) *abstracting the model* and (c) *generalising a result*. In a nutshell, respectively, applying the model to a new system (specifically) means showing that a model may be modified in such a way that it explains or predicts phenomena in new systems; abstracting the model means that the modeller chooses to account for a larger number of systems by describing them less precisely; and generalising a result means showing that a given result holds for a larger number of systems subsumed by the model's described system. Increasing the number of systems subsumed by a model's described system can be achieved both by abstracting the model and by generalising the result. Only the latter is intrinsically epistemically beneficial, however, because abstracting the model implies a loss in the ability to describe some properties of the target. In contrast, when results are generalised, the model retains the expressive power to describe the result as well as all the other properties of the target.¹ On the other hand, if abstracting a model enables proving new important results about targets characterised in a sparser way, those results may be important contributions. It is just that unlike generalising a result, abstracting a model and applying a model to a new system do not provide intrinsic epistemic benefits in the sense of justifiably increasing the confidence in the truth of the model or its result.

Although the philosophical contributions to generality in economics (Hamminga, 1983; Walliser 1994; Rol, 2008) provide some useful distinctions and concepts, expressing the characteristics of epistemically beneficial generalisations specifically requires developing a new conceptual framework. Unfortunately, generalisation is conceptually surprisingly complex, thus I need to give specific interpretations to some commonly used concepts such as 'target', 'system' and 'applies', which are more precise and in some cases different from the usual ones. Most of the paper is thus devoted to formulating a notion of generalisation that allows for expressing when and how generalisations can yield epistemic benefits.

The Dixit-Stiglitz (1977) model of monopolistic competition will be used as a concrete example of using the philosophical concepts introduced. Although the

¹ Abstracting the model and generalising results thus differ with respect to 'expressive power'. The differences are complex, and they cannot be dealt with in this paper due to limitations of space. I develop the notions of expressive power and of abstraction in a companion paper 'The epistemic benefits of generalisation in modelling II: expressive power and abstraction' that focusses on the details of model descriptions in generalisations.

model I discuss is from economics, the conclusions and the analysis in this paper are intended to be applicable to any science.

The paper culminates in two examples of epistemically beneficial generalisations (from Krugman, 1980, and Wald, 1951). The characteristics of these examples determine the desiderata for the philosophical concepts needed to account for them: (1) Given that some but not all generalisations yield epistemic benefits, one needs notions of model applicability that are able to distinguish between the different kinds of generalisation. (2) Given that epistemically beneficial generalisations demonstrate how a model result applies to a larger set of possible systems, the notion of applicability must be able to capture changes in the set of possible systems to which the model result applies. (3) Given that epistemically beneficial generalisations arise from changes in model descriptions that describe the same target and the same model result, the notion of applicability must be able to recognise such changes as generalisations.

With these desiderata in mind, the structure of the paper is the following. Section 2 distinguishes between targets and systems and shows how they may enter into proper subset relations. Section 3 discusses how Michael Weisberg's account of model applicability needs to be modified in order to be able to distinguish between different generalisations. It is argued that at least one notion of model applicability needs to be defined as a relationship between model descriptions and systems. The notion of *described systems* is then distinguished from *possible systems*. Section 4 distinguishes between generic and specific applicability and Sect. 5 between three different kinds of generalisation using these notions of applicability. I give a rough outline of the Dixit-Stiglitz (1977) model of Monopolistic Competition (henceforth abbreviated as MC) in Sect. 6 and illustrate the use of generic model applicability with it. Section 7 describes an example of applying a model to a new system, and shows how one model may explain more phenomena than another while applying to fewer systems. In other words, here a model is generalised according to specific applicability but made less general according to generic applicability. Section 8 illustrates epistemically beneficial generalisations with two examples.

2 Targets, systems and conceptualised systems

Matthewson and Weisberg (2009) discuss model generality in terms of the number of phenomena to which a model relates, and the number of targets to which it applies. Weisberg (2013, pp. 90–93) explains that targets are abstractions over phenomena in that modellers are not concerned about accounting for all aspects of phenomena or systems of interest. When they choose an 'intended scope' for a model, they concentrate on describing some properties and abstract away others. 'This yields a target system, a subset of the total state of the system' (2013, pp. 90–101). A target is thus a matter of the modellers' intentions.

In cases of generalised modelling an intersection of sets of features is an informative generalized target (p. 117). For example, economists may be interested in modelling the shared properties of all monopolistically competitive

markets. The shared properties constitute the generalised target. Here it must include *product differentiation* and the concomitant *monopoly power*.

The target determines which aspects of some system the modeller wants to account for. A *system* is interpreted as encompassing all its properties. Unlike a target, it is not an abstraction. There are various ways of partitioning the world and speaking about different systems in it. Some of the partitions are not related to modelling or targets (unlike in Elliott-Graves, 2018). One might meaningfully talk about international and domestic trade without specifying in which of their aspects a modeller might be interested, for example. Let us refer to the results of such partitioning as *conceptualised systems*. Such mental representations of systems are also abstract: when humans think about various systems, their conceptions of them do not include all their properties. The reason why it is necessary to introduce conceptualised systems is that they enter into *proper subset relations* that are relevant to important kinds of generalisation. Let $S = \{p_1, \dots\}$ refer to a conceptualised system S that has *property* p_1 and some other unspecified properties. For example, let p_3 stand for product differentiation, and p_1 for a mark-up. Such notation allows us to represent various kinds of conceptualised systems:

$$S_1 = \{p_1, p_3, \dots\}$$

$$S_2 = \{p_1, p_3, p_4, \dots\}$$

$$S_3 = \{p_1, p_4, p_5, \dots\}$$

$$S_4 = \{p_1, p_2, p_3, p_5, \dots\}$$

$$S_5 = \{p_1, p_3, \sim p_4, \dots\}$$

The expression ‘ $\dots$ ’ refers to the idea that real systems have many properties that mental representations ignore. Conceptualised systems enter into proper subset relations with each other because they ignore many details of the systems. Here, for example, S_2 , S_4 and S_5 are proper subsets of S_1 . Let us give the properties the following interpretations:

p_1 : monopoly power.

p_2 : a market for hookahs.

p_3 : product differentiation.

p_4 : an international market.

p_5 : a market in which the total value of the product is less than one billion dollars per year.

For example, S_1 is then the conceptualised system of all MC markets, and S_2 is an international MC market, and so on. Conceptualised systems should be called system-kinds rather than systems: each conceptualised system may contain several actual systems that are similar only in terms of having the indicated properties in common. For example, S_2 refers to a market for ice-hockey sticks and a market for shampoo, and many others.

A target *subsumes* one or more systems, such that the cardinality of the subsumed systems is given by the number of systems that share the properties that define a *target*. A target thus picks out selected properties from conceptualised systems, viz. those for which the modeller intends to account with a model. Let us write $T = \{p_1\}$ to indicate that a modeller intends the model to represent property p_1 . If the modeller only wants to account for monopoly power (p_1), target T subsumes all the systems S_1 to S_5 in virtue of the fact that they all have property p_1 . If another, more specific

target T_1 is defined by $T_1 = \{p_1, p_3\}$, then T_1 , the MC market, subsumes systems S_1 , S_2 , S_4 , and S_5 but not S_3 because S_3 does not have property p_3 . The formulation $T_1(S_1, S_2, S_4, S_5)$ describes this subsumption relation.

The distinction between systems and targets, and the subsumption relations will be needed for defining the generic notion of applicability. Precise notions of applicability are needed because generality is defined in terms of applicability, and different notions of applicability are used in different kinds of generalisation.

3 p-generality, applicability and described sets

Michael Weisberg has provided an account of model applicability. Although I will use some parts of his account in accounting for some kinds of generalisation, I cannot do so without considerably modifying the account. Given that Weisberg defines generality as a property of the relationship between models (interpreted as abstract structures) and target systems, generality and applicability are not a properties pertaining to the relationship between model descriptions and (target) systems in his account. We will see that differences in model descriptions of the same conceptualised system are important for epistemically beneficial generalisations. We will thus need a notion of applicability and generality that concerns the relationship between model descriptions and systems.

Weisberg defines applicability as follows (2004, p. 1076): ‘A model applies to a target system when it accurately describes the structure and dynamics of the system according to the standards set by the model builder or model user’. Moreover, ‘A model can be successfully applied to a target when it fits the target ... Modelers’ fidelity criteria specify which properties must fit and to what degree they must fit’ (2013, p. 93). The intended scope defines the target (2013, p. 40). The choice of a feature set as equivalent to the choice of the modeller’s intended scope, which is connected to the modeller’s choice of target (see also Weisberg, 2012, 2015; Parker, 2015).

Matthewson and Weisberg (2009) distinguish between *a-generality*, which is the number of actual targets to which a model applies, and *p-generality*, which is the number of (logically) possible targets to which it applies (see also Weisberg, 2004, 2007; Matthewson, 2011). A modeller may intend to account for properties of a system that is known not to exist in actuality, such as xDNA, or a system that is known to be impossible, such as a perpetual-motion machine (see Weisberg, 2013, pp. 121–129). In these particular cases the distinction between targets and systems makes little difference, and it does indeed make sense to talk about ‘possible targets’ because it is clear that the modellers’ intention has been to describe systems that are known not to exist.

However, various authors refer to ‘possible targets’ when the modellers primarily intend to account for some actual systems. Furthermore, some of them argue that the benefits of generalisation concern possible rather than actual systems (Strevens, 2004, 2008; Potochnik, 2017; Weisberg, 2013, p. 110). The idea is that modellers may generalise their models by showing that there is a larger number of possible systems to which they apply, but the number of actual systems may remain the

same. Orzack (2012) interprets the difference between *a*- and *p*-generality in terms of applicability: increasing the set of ‘possible systems’ means increasing the set of systems to which a model potentially applies. The epistemic benefits of generalisation are related to ‘possible systems’, but not any kind of possible system will do, and it is possible that the benefits also concern actual systems.

Weisberg recognises the existence of situations in which a ‘model has more concrete features than the abstracted target’ (2013, p. 117). Modellers typically need to make overly specific assumptions, many of which are known to be false. Modellers are able to generate the desired features only by assuming the falsities. Weisberg posits that if some features of the model cannot be part of an abstracted target, the modellers should specify that these features remain outside the model’s scope. Suppose that modellers do this, and that they subsequently develop a model with the capacity to correct for how such features are represented, thereby generalising it. Now the ‘model can be constructed at the appropriate level of abstraction’ (ibid.). The target and the actual system remain fixed, but the model descriptions characterising them constitute different possible systems, and the more general model may give a better description of such systems in the sense of making fewer idealisations.

As I will demonstrate, generalising the model in this kind of situation typically yields epistemic benefits, but is this a situation in which a model has been generalised according to Weisberg’s account? If the new model expands the set of possible systems to which it applies, then there could be a generalisation in terms of such possible systems. One prominent interpretation of the notion of possible systems is that they consist in whatever the model descriptions specify. For example, in Giere’s (1988, p. 83) account the abstract object picked out by the model descriptions has all and only those properties that the model descriptions specify (see Thomson-Jones, 2010 for a discussion).

Note that a model cannot fail to apply to a possible system if the possible system is interpreted as a system defined by the model descriptions. If model descriptions define the properties of a possible system, then every model obviously applies to any possible system that it defines. A definition of model applicability can be based on counting the number of possible systems to which the model applies only if a given model description may apply to a *set of* possible systems, and if there is a way of comparing the number of such systems in different model descriptions. In what follows, I will call the set of systems given by a model’s model descriptions its *described system*, S_D . Described systems are different from conceptualised, possible, actual as well as from ‘model systems’.

I will now argue that the different described systems may subsume different numbers of possible systems, and this gives a way of comparing different models with respect to generality. But how exactly is the number of possible systems subsumed by a given described set determined?

Weisberg claims that a given model description may give rise to several models, and argues that one can increase a model’s generality by increasing the number of models that the model descriptions pick out. Given that he does not define generality in terms of a relationship between model descriptions and systems, and given that model descriptions give rise to several possible systems (and ‘models’), I interpret him as arguing that the set of possible systems is not uniquely determined by the

model's described system (see 2013, pp. 34–35). If it is not, then something else determines the set of possible systems to which a model applies. One option is that the possible systems have some properties that are not included in the model descriptions. If this is his view, it is similar to that of the fictionalists. 'Model systems' is a term used by some proponents of the fictional view of modelling (e.g., Frigg, 2010; Godfrey-Smith, 2006). The model system is distinguished from the model description that specifies it in that the former is imagined, and it is the imagined model system that represents real-world phenomena (see also Levy, 2015; Knuuttila, 2017). In other words, model systems contain imagined properties not included in the model descriptions.

The set of possible systems to which a model applies will then include all the possible ways of adding and subtracting, in one's imagination, properties and relations within the constraints provided by the described set. However, there will be an infinite number of such systems if all logically possible systems are admitted, and far too many if all physically possible systems are admitted. This is why it is not acceptable to let any of these imagined properties determine the (generic) applicability and thereby the generality of a model.

I take this to be a slightly different infinity problem than that identified by Strevens (2004) and Lewis and Belanger (2015). They argue that models with real-valued variables have an uncountably infinite number of possible systems. One can solve both infinity problems by arguing that although one cannot count possible systems, it is sufficient for generality comparisons if the sets of possible systems subsumed by described sets enter into proper subset relations.

Lewis and Belanger (2015) provide an extension to cases in which the proper subset criterion is clearly too strict by using measure theory. In what follows, I will take differences with respect to applicability concerning possible systems to be defined by the proper subset relations between different described sets. I do not see why this account could not be extended with measure theory, but I will not make any attempt to provide such an extended account here.

Weisberg (2004) provides an account that involves proper subsets, but with different terminology.² The proper subset solution to the infinity problem is not available, however, unless the set of possible systems to which a model applies is strictly determined by the model descriptions. If it is not, and the imagined

² Weisberg (2004) argues that the precision of model descriptions determine the number of 'models' picked out by those descriptions. If these 'models' then correspond to the possible systems, Weisberg's solution is consistent with what I say in the text. The 'models' that Weisberg discusses here are abstract structures or points in state space (see e.g., 2013, p. 42). Weisberg shows that generality may increase if the model descriptions pick out a proper superset of models (qua abstract structures), and the models must then apply to a larger number of possible systems. I do not need to assume that the possible systems correspond to models qua abstract structures or to model systems, nor that such abstract structures or model systems even exist. My account may thus be more palatable to philosophers who think that such things either do not exist or are epistemically suspect or unnecessary for understanding modelling (e.g., Levy, 2015; Odenbaugh, 2018). However, I do not need to deny the existence of abstract structures either, as long as at least some kinds of generality and applicability are conceptualised as having to do with the relationship between model descriptions and systems rather than models qua abstract structures and systems.

properties play a part as well, then there is no guarantee that modellers will ever come to an agreement about the set of possible systems to which a model applies.

In Weisberg's account, applicability of a model is determined by fidelity criteria (2013, pp. 42–43, p. 93). Re-writing Weinberg's definition in terms of possible systems indicates why this may be a problem: 'A model applies to a possible system when it accurately describes the structure and dynamics of the system according to the standards set by the model builder or model user.' The problem is that it only makes sense to talk about fidelity criteria with respect to actual but not possible systems. A more serious problem is that changes in the proper subset relations that define the generality of a model cannot be changed by fidelity criteria because such subset relations are objective features of the model descriptions rather than changeable through altering the fidelity criteria. It is the model descriptions that allow for demonstrating subset relations among sets of possible systems, and they determine the (generic) applicability of a model.

To put it differently, applicability and thus generality are determined in two different ways in Weisberg's account: through the precision of model descriptions and through the fidelity criteria. In effect, then, Weisberg has been discussing two different notions of model applicability without apparently realising it. Instead of trying to reconcile them, I will argue that there are two different notions of applicability that give rise to three different notions of generality. We will then see in Sect. 7 that the two notions of applicability may give conflicting judgments concerning whether a model-modification counts as an increase or a decrease in generality.

4 Generic and specific applicability

A model can be said to apply to a conceptualised system in two different senses, *generically* or *specifically*. In the case of generic applicability, the properties in the target may be shared by a large number of different systems, and a model applies to a given (target) system just in case it is able to describe the properties in its target. If it does, the target subsumes the system along with the other systems. This kind of applicability fits well with what Weisberg calls 'generalised modelling', and the generality of models is evaluated by proper subset relations into which the described sets of two or more models enter.

A model applies to a system specifically if it is able to describe some system-specific properties. This notion of applicability is often used when an existing model is modified so as to be applied in some new circumstances. In such cases modellers often use fidelity criteria in deciding whether the model applies to a given system.

The notions of applicability are formulated in terms of whether a model applies to a system rather than to a target. The target determines which properties the model must be able to represent. The definitions are success-based in the sense that the model must indeed be able to represent the properties defined by the target.

Generic model applicability (GAPP) requires that

(GAPP) A model M applies generically to a system S if it successfully represents the features that define the target T , and system S has the properties that define the target T .

For example, given that the target of models of monopolistic competition consists of product differentiation and monopoly power, any model in which the model descriptions are able to represent these properties applies generically to such markets. Krugman's (1980) model and the Dixit-Stiglitz model (1977) apply generically to such systems.

Suppose now that M already applies to S according to GAPP and S' is a proper subset of S . Specific applicability (SAPP) requires that.

(SAPP) A model M' applies specifically to S' if M' is able to describe some S' -specific properties of S' .

SAPP is also used when a model is said to apply a pre-existing model to a new conceptualised system.

(SAPP) Model M' applies model M specifically to system S' if M' is able to describe some S' -specific properties of S' in virtue of using the derivational and conceptual resources of M .

For example, Krugman's (1980) model applies the Dixit-Stiglitz model M_{DS} specifically to the conceptualised system of international markets in virtue of using resources from M_{DS} in describing some characteristics specific to international markets. The distinction allows for making sense of cases in which a model applies to a system S' but does not do so specifically. The Dixit-Stiglitz (1977) model thus applies generically to international markets, but it does not apply specifically to them.

The notion of *applying a model* M specifically to a system S' presupposes that model M already applies to some set of conceptualised systems S according to GAPP. S' must already be among the systems subsumed by the target T , and hence a proper subset of S . If it were not, it would be impossible to use the conceptual and derivational resources of M in applying it specifically (with M') to S' . To put it differently, one cannot use the conceptual and derivational resources of M in accounting for S' with M' unless S and S' share some properties, namely, those properties that the conceptual and derivational resources of the two models are able to describe. It follows that the targets T and T' must share some properties. On the other hand, the targets cannot be identical because T' must include some S' -specific properties that are not in T . The model-version M' that is responsible for applying model M specifically to S' cannot be identical with M . One is justified in saying that model M' applies specifically or generalises M only if the two models are sufficiently similar so that the former can be said to use the conceptual and derivational resources of the latter.

5 Three kinds of generalisation

We are now almost in a position to use the two notions of applicability for defining three different kinds of generality. All we need is a notion of a *model result*. What does it mean for a model M to successfully represent the properties of a system? It means that the properties of a system, say $S_1, (p_1, p_3)$, are being represented by *model descriptions* in M , which may consist of *assumptions* A_i , *model implications* or *model results*.

A *model result*, R , is a proposition derivable from a model that is thought to be epistemically important within the appropriate scientific community. Given that all assumptions are propositions derivable from the model it is necessary to distinguish them from model results, which must be derivable only from some combination of the assumptions rather than any assumption or functional form alone. Since for any two assumptions A_1 and A_2 , $(A_1 \& A_2) \vdash A_1, A_2$, R is a result of a model $(A_1 \& A_2)$ only if $(A_1 \& A_2) \vdash R$ and $A_1 \not\vdash R$ and $A_2 \not\vdash R$. This requirement is too weak, however, because functional forms in mathematical models typically entail several assumptions. If a functional form F alone entails a result, such a result is indistinguishable from an assumption: If $F = (A_1 \& A_2)$, $F \vdash A_1, A_2$ but it would be correct to call A_1 and A_2 *model results* only if the model consisted of nothing else but this one functional form. What is missing here is the idea that model results must be generated by the joint operation of different model components. One can define model results by requiring that there are at least two functional forms F_1 and F_2 such that $(F_1 \& F_2) \vdash R$, $F_1 \not\vdash R$ and $F_2 \not\vdash R$.

Model implications are like model results except that they have been derived in previous modelling exercises. A model result of an earlier model may thus become part of a target of a new model if a modeller requires that such model implications must be derivable with the new model.

Suppose that model M applies generically to (a set of systems) S according to GAPP in virtue of the fact that its described system S_D successfully represents the properties in target T . Model M' can then be taken to generalise M if.

- (a) M' applies M specifically to another actual system S' and another target T' , or if
- (b) M' applies to actual system S (or to actual system S^* such that S is a proper subset of S^*), such that S_D is a proper subset of its described system $S_{D'}$, and T is a proper subset of T' , or if
- (c) M' applies to actual system S , such that S_D is a proper subset of its described system $S_{D'}$, and the target remains the same.

(b) and (c) rely on generic and (a) on specific applicability (and M' applies M specifically in a).

These cases correspond to.

- (a) *applying a model (specifically) into a new system,*
- (b) *abstracting the model and*

(c) *generalising the model result.*

(b) and (c) are based on increasing the number of systems subsumed by a model's described system. Abstracting the model (b) means changing the target in such a way that it subsumes a larger number of systems. Targets thus also enter into proper subset relations. The difference between (b) and (c) is that generalising a result (c) requires that the target remains the same, while abstracting the model means stripping properties away from the target. When the model is abstracted, T is a proper subset of T' because M' does not depict some properties in T . It is possible to explain the difference between abstracting the model and generalising results in considerably more detail, but limitations of space imply that this requires another paper.³

Generalising a model always presupposes that there already is a model that applies to some actual system according to GAPP. If the modeller generalises the model but keeps the target the same, then, again according to GAPP, the model applies to the same actual systems S as before. In other words, a model modification that generalises the model result (c) cannot increase the number of actual systems to which the *model* applies because the sameness of the target implies that the model must have already applied to the actual systems that the target subsumed. Changing the described set cannot change the actual systems to which the model applies as long as the target remains the same. This is why, according to the definition (b), abstracting the model may change the number of actual systems to which the model applies, but (c) generalising a result cannot change it because such generalisations keep the target fixed. We will see later, however, that generalising the result may increase the set of actual systems to which the *model result* applies.

It is easier to grasp such claims by considering concrete examples. I will introduce the Dixit-Stiglitz model in the next section, and then look at examples of generalising a model by (a) applying it to a new system (Sect. 7) and of (c) generalising the result (Sect. 8).

6 The Dixit-Stiglitz model of monopolistic competition

The Dixit-Stiglitz (henceforth DS) model is particularly suited to my purposes because it has been widely applied in circumstances that differ considerably from those in the original (1977) model, which was intended as a contribution to the theory of industrial organisation. It has been applied to several different phenomena in various fields such as international trade, geographical economics and macroeconomics. I will discuss some of these applications later.

Perfectly competitive markets have 'homogeneous' goods: each firm produces an identical good, and such goods are said to be 'perfect substitutes' to each other. Consequently, if a firm charges a lower price for the good than others, given various further assumptions about the markets, every consumer is assumed to buy the good from this firm only.

³ See 'The epistemic benefits of generalisation in modelling II: expressive power and abstraction'.

A *monopolistically competitive* (MC) market is characterised by *product differentiation*: each firm produces a single good that is a close (but not perfect) substitute for other goods in the sector. Different firms provide slightly different kinds of shampoo, for example: for dry hair, for dyed hair, sports shampoo, anti-dandruff shampoo and so on. Such product differentiation yields some *monopoly power* to these firms, which are able to charge a price that exceeds marginal costs. The difference between the marginal cost and the price is called the *mark-up*.

Edward Chamberlin first developed the theory of monopolistic competition in the 1930s. One of the perceived problems of his account was that it turned out to be almost impossible to determine the exact scope of MC markets. Rather than solving this problem, Dixit and Stiglitz (1977) purported to describe such markets by introducing the sub-utility function $V(x_1, \dots, x_n)$, which is symmetrical in its arguments x_i . Their ‘solution’ consisted in introducing a false idealisation, but one that allowed them to present a model that integrated monopolistic with perfect competition. Their main contribution was to present a tractable *general equilibrium* model (Neary, 2004), which was capable of treating all markets simultaneously. The *model result* was that the equilibrium number of varieties in the MC industry is optimal.

The basic idea is that firms and consumers determine the allocation of goods first between the perfectly competitive and the MC sectors, and then within the MC sector. Free entry then determines the number of firms in the MC sector. Let n denote the number of products or firms in the group, x the output of each such firm, and p the price of each product in the group relative to a composite good (an aggregate of all other goods in the economy). Let $C(x)$ denote each firm’s cost of production, and let $\varepsilon(p, n)$ denote the elasticity of demand perceived by each firm. The important equilibrium condition in the model is

$$p \left[1 - \frac{1}{\varepsilon(p, n)} \right] = C'(x) \quad (\text{Eq})$$

Perfectly competitive markets can be obtained as a limit case of the DS model, given the suitable choice of parameter values: $\varepsilon(p, n) = \infty$. There is a positive mark-up if $\left[1 - \frac{1}{\varepsilon(p, n)} < 1 \right]$, in other words, if $\varepsilon(p, n) < \infty$. This is why consumers’ preference for variety and firms’ monopoly power are really two sides of the same coin in the original DS model: preference for variety allows for deriving $\varepsilon(p, n) < \infty$, and this leads to a positive mark-up, as indicated in equation (Eq).

Dixit and Stiglitz consider three restrictions on the utility function U , and analyse three cases in which they relax one restriction at a time: symmetry of V in the x_i , a Constant Elasticity of Substitution (CES) functional form for the sub-utility function $V \left(V = \left(\sum_{i=1}^n x_i^\rho \right)^{1/\rho} \right)$, and a Cobb–Douglas form for U itself ($x_0^{1-\mu} V^\mu$). A constant elasticity of substitution means that consumers’ willingness to substitute x_1 with x_2 does not depend on the consumption of $x_3, x_4, \dots$. Incorporating all the restrictions produces the utility function

$$U = x_0^{1-\mu} V^\mu, \quad V = \left(\sum_{i=1}^n x_i^\rho \right)^{1/\rho} \quad (\text{DSP})$$

here, $1-\mu$ is the share of nominal income spent on the composite good x_0 , and ρ measures the substitutability among product varieties. When economists talk about ‘Dixit-Stiglitz Preferences’, they usually refer to (DSP). Using a ‘love of variety’ utility function such as the CES allows modellers to derive $\varepsilon(p,n) < \infty$. In other words, consumers’ preference for variety creates monopolistic markets in which firms face a less than perfectly elastic demand. Almost all applications of the model include all three restrictions such that (DSP) is assumed to hold (Neary, 2004). With the main concepts of the model to hand, it is possible to see how it has been modified, applied and generalised.

Let us now see how GAPP is to be interpreted by considering the example of monopolistic competition. Implications such as a positive mark-up in a model of monopolistic competition are successfully represented just in case they are derived with a model. Thus, models of monopolistic competition never include the *assumption* of a positive mark-up, and the model describes this property of real markets only when such a mark-up is successfully derived from the assumptions. This is why model descriptions must be taken to include model implications.

If a modeller wants to study monopolistic competition, the target consists of product differentiation p_1 and monopoly power p_3 : $T_1 = \{p_1, p_3\}$, and the model M_1 applies to systems S_1 , S_2 , S_4 , and S_5 if it successfully describes properties p_1 and p_3 . If it does successfully describe those properties, let us write the *descriptions of properties* p_1 and p_3 , d_1 and d_3 , on the right-hand side of the turnstile ‘ $\vdash$ ’: $M_1 \vdash d_1$, d_3 . A described property d_1 corresponds to the real property p_1 , d_3 to p_3 , and so on. In such a case, let us say that a model has the *expressive power* to describe properties p_1 and p_3 .

The modeller’s intentions can now be placed on the left-hand side of the symbol $\vdash$. The expression.

$$M_1(T_1(S_1)) \vdash d_1, d_3$$

means that the modeller intends model M_1 to apply to conceptualised system S_1 that has the properties defined by target T_1 , and that it does apply to this system in virtue of successfully representing the properties that define the target. For example, a modeller intending to describe a mark-up (p_1) will do so successfully if one can derive $d_1 = \varepsilon(p,n) < \infty$ from the model:

$$M_1(T_1(S_1)) \vdash d_1 \text{ because } M_1(T_1) \vdash \varepsilon(p, n) < \infty$$

We can now state that model M_1 applies to MC markets (S_1) in virtue of the fact that it successfully describes a mark-up (p_1) and product differentiation (p_3) because it entails descriptions d_1 and d_3 as model implications.

$$M_1(T_1(S_1)) \vdash d_1, d_3 \text{ and } S_1 = \{p_1, p_3, \dots\}$$

Note that representing S_1 entails representing S_2 , S_4 , and S_5 as well, because these are all proper subsets of S_1 . Hence,

$$M_1(T_1(S_1, S_2, S_4, S_5)).$$

The requirement of being able to represent the features of the target with *model descriptions* makes stringent demands. A model could fail to apply to a system

generically if its described system does not contain the features that define the target. Consider, for example, the abstract functional form for utility that Stiglitz and Dixit (1993) introduced to describe the bare bones of their model.

$$u(x_0, V) = U(x_0, V(x_1, \dots, x_n)) \quad (\text{SE})$$

here utility depends on consumption of the composite good x_0 and on differentiated products x_i . There is a sense in which this functional form is intended to apply to *all* markets: it provides a general equilibrium account. On the other hand, if one requires (and economists surely do require) that MC markets must be modelled so as to allow for representing product differentiation and market power, then this functional form does not apply to such markets because it cannot represent these features. In particular, it does not allow for deriving any particular value for $\varepsilon(p, n)$ because it is too abstract. As a result, it does not have the expressive power to describe d_1 and d_3 . Mere intention to represent is thus not sufficient for generic applicability. Descriptions of these features of monopolistic competition can be derived from a model only if V is further specified as a preference-for-variety utility function such as the CES function: $V = (\sum_{i=1}^n x_i^\rho)^{1/\rho}$.

GAPP is rather permissive in that it does not specify *how well* a model ought to apply to a conceptualised system. Whatever criteria modellers use for deciding the features of a target, if the model descriptions represent those features, the model applies generically to any system that has them. This notion of applicability does not rule out model descriptions that are too specific and/or highly idealised. They may mis-describe the systems subsumed by the target, or provide a description that truthfully applies only to a subset of systems that have the characteristics that define the target. The fact that a model is idealised does not thus rule out generic applicability. For example, a model M may successfully apply to system $S_4 = \{p_1, p_2, p_3, p_5, \dots\}$ if the target is $T_1 = \{p_1, p_3\}$ and if.

$M(T_1) \vdash d_1, d_2, d_3$.

However, models M_2 and M_3 also apply to S_4 if.

$M_2(T_1) \vdash d_1, d_2, d_3, d_4, d_5$ and if

$M_3(T_1) \vdash d_1, d_2, d_3, d'_4, d'_5$,

where d'_4 is incompatible with p_4 , and d'_5 is incompatible with p_5 .

7 Applying a model to a new system

I will now provide an example of generalising a model by applying it to a new system (a). I will show that a given model-modification may simultaneously increase the number of phenomena to which a model applies according to SAPP, and decrease the number of conceptualised systems to which it applies according to GAPP. This will happen if a model relates to more phenomena than another model, but yet applies to fewer systems than the latter.

Various philosophers talk rather loosely about ‘target phenomena’. Unfortunately, such an expression may give the impression that a phenomenon always occurs in

a corresponding system. In the context of monopolistic competition, phenomena include mark-up pricing, consumer preference for variety and the associated product differentiation by firms. These phenomena are constitutive of the conceptualised ‘MC market’ system. A given phenomenon may occur in several different conceptualised systems and, obviously, a given system may give rise to several phenomena.

Consider Krugman’s (1980) application of the DS model to international trade. Krugman’s model could be said to apply the DS model specifically to a new system on the grounds that it uses the derivational and conceptual resources of the DS model in making sense of a phenomenon in this new system.

Krugman shows why countries have an incentive to trade similar goods ($=R_1$). This is a system-specific phenomenon that had not previously been explained, and here the phenomenon corresponds to the model result. Krugman (a) *increases the number of phenomena and the number of conceptualised systems* to which the DS model applies specifically in that his model uses its conceptual and derivational resources to explain a new phenomenon.

Krugman uses the following CES function:

$$U = \left(\sum_{i=0}^n x_i^\rho \right)^{1/\rho} \quad (\text{K})$$

Recall that the most commonly used version of Dixit-Stiglitz preferences looked like this:

$$U = x_0^{1-\mu} V^\mu, V = \left(\sum_{i=1}^n x_i^\rho \right)^{1/\rho} \quad (\text{DSP})$$

(K) is a special case of (DSP) when $\mu = 1$. In other words, Krugman’s model describes the international market as if it only had a MC industry. Krugman does not provide a genuine general-equilibrium model because (K) does not represent the relationships between a MC industry and the other industries. It applies generically to fewer kinds of systems than the DS model and it is less general in this sense.

Krugman’s explanandum phenomenon was the existence of international trade in similar commodities. The DS model allows for studying the general equilibrium consequences of different relative shares or, as in the more abstract specification $U(x_0, V = (\sum_{i=1}^n x_i^\rho)^{1/\rho})$, represents situations in which the MC market is effectively independent of the perfectly competitive markets.

Let p_r denote such general equilibrium consequences. It is now possible to describe the target of the original DS model more precisely as $T_{DS}(p_1, p_3, p_r)$ and Krugman’s as $T_K(p_1, p_3, p_4)$. The systems subsumed under T_K are a proper subset of those subsumed under T_{DS} for two reasons. T_K includes neither perfectly competitive nor domestic markets. Krugman’s model shows that firms have an incentive to trade similar goods (R_1) in a proper subset of the set of all MC markets (S_1), namely the international ones (S_2). Thus, whereas Krugman’s model applies (specifically) to a new phenomenon, the new phenomenon concerns a proper subset of the systems that the DS model subsumes.

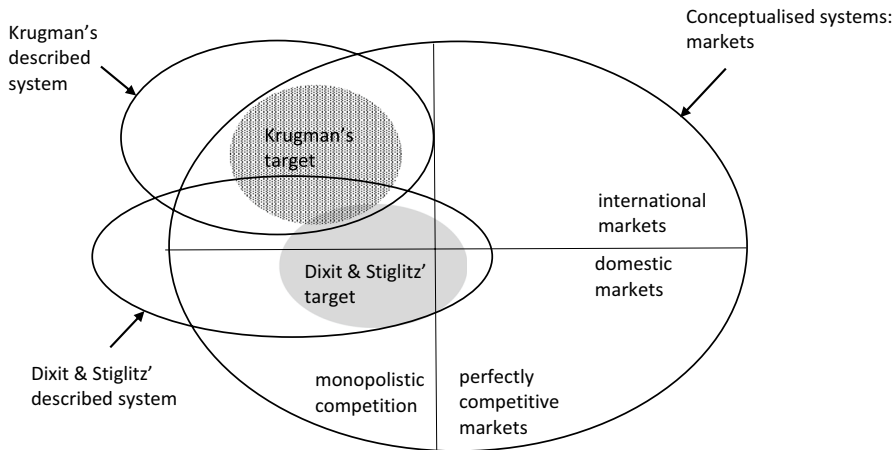


Fig. 1 A comparison of systems and targets

Figure 1 illustrates these relationships. The large oval depicts the properties of all markets. The upper half depicts the properties of international markets, and the lower half the domestic ones. Similarly, the left side depicts MC markets, and the right side perfectly competitive ones. Since Krugman aims to explain properties specific to international markets, his target covers more of the international markets, and his described system covers more properties of international markets with monopolistic competition. Yet it does not apply to perfectly competitive markets even generically because it does not have enough expressive power to describe them.

Krugman's model is often taken to apply the DS model specifically to international trade because the latter does not have the expressive power to imply anything specific about international markets. On the other hand, the DS model already applies to international markets according to GAPP. Some international markets are characterised by monopolistic competition, and the DS model applies to those markets in virtue of the fact that they are systems subsumed by the set of all MC markets. To put it differently, since the DS model describes monopolistic competition in all markets, it thereby also describes it in international markets.

Krugman's model is a typical application of a model which is more general in terms of generic applicability. A specific model is built by adding assumptions to it. However, the general model is not sufficient in itself for deriving such results because it simply does not reveal anything about the details of various systems. One possible way of characterising the notion of applying a model specifically to a new system (a) in this case is to say that the DS model is a *general model*, and that Krugman's model applies but does not generalise it. Then Krugman's model would not count as a generalisation of the DS model, but it would count as a *demonstration that* the DS model is indeed general in the sense of being widely applicable.

However, scientists may talk about 'applying a model (specifically) to a new system' even when the model applied and the applying model are at the same level of generality. For example, a model developed for a specific institution or organism

may be applied (specifically) to other institutions or organisms of the same kind. In a case like this, the exact identity of the ‘general model’ is more diffuse because the model is never formulated in its most general form. Consider, for example, McKelvey and Ordeshook’s (1972) model of voting in the so-called plurality rule. This model formulates ‘expected gain’ expressions for various candidates that are based on the idea that voters condition their choice on being pivotal. Merrill (1981) formulated similar expressions for the Borda rule and approval voting. Unlike in the case of the Dixit-Stiglitz model, it is impossible to write down the model descriptions for a ‘general model’ that would give rise to these specific applications because the equations for expected gain depend on the details of the voting rule in question.

Yet voting theorists posit that Merrill applied McKelvey and Ordeshook’s model into new conceptualised systems or even that he generalised it. They do so because it is possible to use the conceptual and derivational resources ‘of a model’ in new circumstances even when there is no canonical and general representation of the model. ‘The model’ here refers to a set of model-building principles. Although this may sound vague, it is not necessarily difficult to tell the different models apart. For example, Enelow (1981) provides an expected-utility model and Ordeshook and Palfrey (1988) an expected-gain model of strategic voting for the same voting rule.

When a model is extended in order to derive further results, that is, to account for new phenomena about the *same* conceptualised system, I will use a *script* font for the term *applies* because this kind of model-application is not covered by the term ‘applies specifically’. For example, Lehtinen (2007) *applies* Enelow’s model to study a new phenomenon (the welfare consequences of strategic voting) under the same voting rule, and Lehtinen (2015) generalises the (2007) model and its *results* to concern any number of candidates instead of just three. Lehtinen’s (2007) model studies a different phenomenon and has a different ‘target phenomenon’ than Enelow’s model but Lehtinen (2007) does not generalise Enelow’s model because he merely derives further results from the model in the same conceptualised system. I take this example to indicate that not every *application* of a model counts as a generalisation.

One might challenge my claim that Lehtinen (2007) does not generalise Enelow’s model on the grounds that explaining any new phenomena is sufficient for exhibiting generalisation. Then, I could have just as well called (a) something like ‘explaining new phenomena’. Ultimately, however, arguing about the semantics is not important. What does matter for the purposes of this paper is whether these different notions of applying specifically and *applying* have different epistemic characteristics.

Perhaps (a) has more to do with the static notion of *generality* as being applicable specifically to a large number of conceptualised systems than with the dynamic changes that can be analysed with *generalisation*. In any case, (a) captures the idea that a model can be said to be general in virtue of being *broadly applicable*. Many authors argue that generality is valuable because of its connection to explanatory power (e.g., Strevens, 2004; Weisberg, 2004), and it is commonly understood as increasing the number of phenomena to which a model applies. This kind of generalisation is thus closely related to unification (e.g., Kitcher, 1989). It seems to me that it makes no epistemic difference whether or not one may write down a canonical general model, and whether *application* counts as generalisation. At the very least unificationists do not claim that such differences are epistemically relevant.

In Krugman's case, explaining new phenomena in a new conceptualised system has an epistemic benefit in the form of the new empirical evidence that supports the DS model: Krugman was able to explain the existence of international trade in similar goods, and the evidence concerning this phenomenon can now be counted in favour of the model of monopolistic competition. However, it is not clear which part of the model is confirmed. If one asks whether this new evidence buttresses the truth of Dixit and Stiglitz' model result of an optimal product portfolio, the answer must surely be no. Using the following generalisation of the CES function

$$V = n^{\eta-\lambda} \left(\sum_{i=1}^n x_i^\lambda \right)^{\frac{1}{\lambda}} \quad (\text{CES-G})$$

Brakman and Heijdra (2004) show that there may be too few or too many varieties, depending on the parameter values. Here η parameterizes the market power of producers of differentiated goods, and λ captures the preference-for-diversity effect. Dixit and Stiglitz' original model result thus only holds under the highly restrictive assumption of the CES function, and is likely to be false about most real markets. If Krugman's model has epistemic benefits deriving from unification, they consist in the fact that product differentiation and monopoly power are shown to be relevant in new circumstances. These kinds of benefits are weak compared to generalising results.

To summarise this section, Krugman's model generalises the DS model by applying it specifically to explain a new phenomenon in a new conceptualised system. Yet this phenomenon concerns a proper subset of the conceptualised systems to which the DS model applies generically. In this example, a generalisation in terms of (a) implies a less general model in terms of (b) and (c).

8 Generalising a result

8.1 Increasing the number of systems subsumed by the described system

In this section, I will discuss cases in which a model is generalised in such a way that the model result remains the same. I will first characterise the different systems in such epistemically beneficial generalisations, and then provide examples from real science.

Let us again consider different systems and their properties.

$$S_1 = \{p_1, p_3, \dots\}$$

$$S_2 = \{p_1, p_3, p_4, \dots\}$$

$$S_3 = \{p_1, p_4, p_5, \dots\}$$

$$S_4 = \{p_1, p_2, p_3, p_5, \dots\}$$

$$S_5 = \{p_1, p_3, \sim p_4, \sim p_5, \dots\}$$

The interpretations are identical to the previous ones, except that p_5 now denotes a property that follows from the ordinary CES function but not the generalised one

(CES-G): the substitutability of commodities is proportional to the firms' market power. Consider the following set of models.

$M_{BH}(T_1) \vdash d_1, d_3$	'Brakman and Heijdra's model'
$M_H(T_1) \vdash d_1, d_2, d_3, d_5$	'A model of hookah markets'
$M_K(T_1) \vdash d_1, d_3, d_4, d_5$	'Krugman's model'
$M_{DS}(T_1) \vdash d_1, d_3, d_5$	'Dixit-Stiglitz model'

Insofar as the target is $T_1 = \{p_1, p_3\}$, all of these models apply generically to all the MC markets. Model M_{DS} applies generically to S_5 even though it depicts it as having property p_5 , which it does not have. The model idealises by distorting this property, but this does not prevent it from applying to system S_5 generically if the distorted property is not included in the target.

A model result R is always proven for a given described system: $R(S_D)$. For example, the described system in M_{BH} is (d_1, d_3) . The models and their results can be written as follows:

$M_{BH}(T_1) \vdash R_{BH}(S_{D1}(d_1, d_3))$
$M_H(T_1) \vdash R_H(S_{D2}(d_1, d_2, d_3, d_5))$
$M_K(T_1) \vdash R_{HMR}(S_{D3}(d_1, d_3, d_4, d_5))$
$M_{DS}(T_1) \vdash R_{DS}(S_{D4}(d_1, d_3, d_5))$

S_{D2} , S_{D3} , and S_{D4} must be proper subsets of S_{D1} because they have all the properties that S_{D1} has, and some additional ones. Thus, a result proven for S_{D1} is more general than a result proven for S_{D2} , S_{D3} , and S_{D4} . M_{BH} characterises the described systems on the right level of abstraction in the sense that the model descriptions do not represent the properties that make the conceptualised systems different from each other: it merely represents what is common to them. Any result R derived with M_{BH} correctly describes conceptualised systems (S_1, S_2, S_4, S_5) , whereas M_{DS} correctly describes only system S_4 . However, given that the model descriptions of M_{DS} describe the properties defining the target $T_1(p_1, p_3)$, M_{DS} applies to systems (S_1, S_2, S_4, S_5) just as M_{BH} does.

There is a crucial difference between M_{BH} and M_{DS} : whereas M_{DS} describes the target, it does so by depicting described systems that have properties that the modeller did not intend to represent. The epistemic problem is that the modellers cannot be sure whether the result only holds because of such assumptions unless they successfully generalise the model. In particular, they cannot be sure that the model successfully applies to systems S_1 and S_2 when the result has been formulated in terms of a model that describes system S_4 . Recall that Brakman and Heijdra (2004) generalised the CES function into CES-G, and showed that the result of the optimal production of varieties R_{DS} no longer held: $R_{BH} = \sim R_{DS}$. In other words, they generalised the DS model.

$M_{DS}(T_1) \vdash R_{DS}(S_{D4}(d_1, d_3, d_5))$
by introducing M_{BH} :
$M_{BH}(T_1) \vdash \sim R_{DS}(S_{D1}(d_1, d_3))$

This indicates that the model result R_{DS} crucially depends on d_5 . Furthermore, one can no longer say that result R_{DS} derived from M_4 applies to systems S_1 and S_2 . One must now acknowledge that although model M_{DS} applies to systems S_1, S_2, S_4 and S_5 in virtue of successfully representing target T_1 , result R_{DS} does not.

In contrast, if the modeller successfully generalised model M_H ,

$M_H(T_1) \vdash R_H(S_{D_2}(d_1, d_2, d_3, d_5))$

with M'_H , one could write.

$M'_H(T_1) \vdash R_H(S_{D_2'}(d_1, d_2, d_3)).$

There seems to be some generalisation here. One could surely say that the *result* has been generalised, because the described system S_{D_2} is a proper subset of the described system $S_{D_2'}$. The generalisation would confirm that model result R_2 does not depend on whether the system for which it is proven has property p_5 . The assumptions of model M_4 have thus been ‘relaxed’.

Model M' *generalises the result* of M if the set for which M ’s model-result is proven, $R(S_D)$, is a proper subset of that of M' , $R(S_{D'})$. The idea here is that the result holds for a described system $R(S_D)$, and S_D is often a tiny proper subset of the actual systems S to which the model applies. Indeed, given that S_D is usually idealised in addition to being abstract, there are very few or no actual systems that have the properties in the described systems. In other words, described systems S_D usually contain properties that pertain only to a small subset of the systems S to which the model applies. This is why generalising results may bring epistemic benefits: in showing that the described system subsumes more systems than before, it increases the number of systems in which result R is shown to hold.

It is relatively easy for a model to apply generically to some systems in the sense that it is sufficient if the described system represents the properties defined by the target. As I have shown, this representation may be at the wrong level of abstraction in that it contains extraneous properties that make it look very different from any actual systems. It may seem more natural to posit that one should employ stricter fidelity criteria for the original models, concluding that they should not have been taken to apply to any relevant actual systems in the first place. One could criticise models that one takes to be too abstract or too idealised for not applying to any real systems. These intuitions are quite correct in many cases. Furthermore, the presence of such abstractions and idealisations may be taken to imply that modellers could use fidelity criteria in figuring out about whether a model and its result apply to some system. They may well do so, but if a model result is generalised, this is not a matter of modellers’ judgment.

It is important to realise that the generic applicability of a model to a system is different from the applicability of a *result* to a system. The reason for this is that the *model results* must be formulated in terms of *described systems*, but the described systems are always different from the *actual systems* to which the model may be taken to apply. Modellers compare the described system to the actual system and try to figure out whether the result depends on the idealisations or on the properties that the described system characterises correctly. There may be considerable uncertainty concerning whether the idealisations are likely to be responsible for the model result. They sometimes are, and this explains why generalisations may bring epistemic benefits. This gives modellers the motivation to try to generalise the model in an epistemically beneficial way. If modellers show that some false assumptions do not affect the model result, this increases the confidence that the other parts of the model are sufficient for obtaining the result.

Generalising a model by increasing the number of systems subsumed by the described system does not guarantee that the original *model result* applies to some

actual system. Given that the set of systems subsumed by the target and the set subsumed by the described system are practically always different from each other, it may be that a model applies to a given system but its model result does not, as illustrated by Brakman and Heijdra (2004). It would be wrong to deny on these grounds, however, that the DS model never applied to MC markets.

In order to make sense of this, one must distinguish between properties that define the target, and properties that define the described set and thereby allow using GAPP for determining the generic applicability of a model. Even though model results from earlier modelling efforts may become the targets for later models, the difference between targets and model results is in practice often sufficiently clear. As Jebeile and Barberousse (2016) note, different parts of a model are typically responsible for representing a system and for representing the phenomena in it.

Suppose now that a result is generalised by increasing the set of systems subsumed by the model's described system. When does such a generalisation in terms of *possible* systems also imply that the *model result* applies to a larger number of *actual* systems? As we will see in the next subsection, epistemically beneficial generalisations that increase the number of systems subsumed by the described system typically remove an idealisation in such a way that one can easily see which actual systems would be included in the larger but not the smaller described system. Such generalisations indeed increase the number of actual systems to which the model *result* applies. Thus, although generalisations that bring epistemic benefits are formally a matter of changes in model descriptions that increase the set of possible systems subsumed by the described system, the increase may also involve actual systems. Whether or not they do so depend on the details of the case.

Let us now consider some examples from economics in which there is a clear epistemic benefit from generalising results.

8.2 An example: krugman's generalisation

Note that $M_K(T_1) \vdash R_{HMR}(S_{D3}(d_1, d_3, d_4, d_5))$ corresponds to a rough description of Krugman's (1980) model. The model result is the 'Home Market Result', $HMR =$ 'each country exports the goods in which it has a large domestic market'. Let us write it in some more detail as follows:

$$M_K \vdash HMR(d_{mc}, d_4, d_{CES}, d_s, d_h, d_{size}, \dots).$$

Here, the idea is that model result HMR holds in a described system that has properties $d_{mc} = d_1 \& d_3, d_4, d_{CES}, d_s, d_h$, and d_{size} with the following interpretations:

$d_4 =$ 'is an international market',

$d_{mc} =$ 'is a MC market',

$d_{CES} =$ 'consumers' preferences exhibit constant elasticity of substitution',

$d_s =$ 'the commodities are symmetrically distributed in the monopolistically competitive market',

$d_h =$ 'the utility functions are homothetic',⁴ and.

$d_{size} =$ 'there are two countries that have the same population size'.

⁴ Increasing each good by a constant amount λ in a homothetic utility function does not affect the elasticity of substitution between commodities: $U(x_0, V(\lambda x_1, \lambda x_2, \dots, \lambda x_n)) = U(x_0, \lambda V(x_1, x_2, \dots, x_n))$.

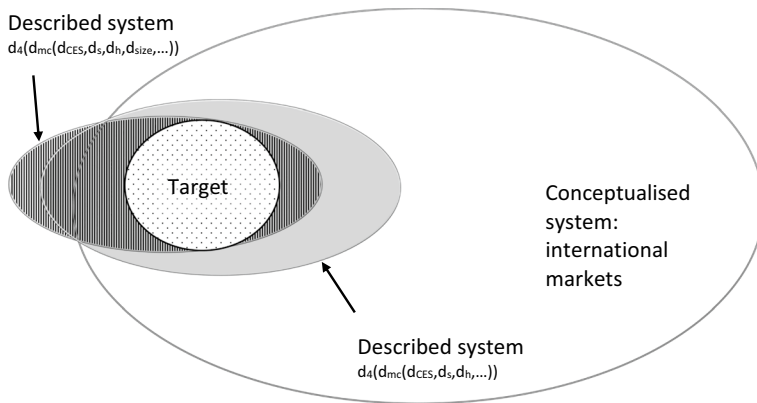


Fig. 2 Described systems in Krugman's generalisation

Given the assumption in Krugman's model that all international trade is characterised by monopolistic competition, one could write the described system as follows:

$$M_K \vdash \text{HMR}(d_4(d_{mc}(d_{CES}, d_s, d_h, d_{size}, \dots))).$$

This formulation indicates that model result HMR concerns a description of international trade (d_4) that is characterised by monopolistic competition (d_{mc}) and no other market forms, and that the thus characterised system is further restricted by d_{CES} , d_s , d_h , and d_{size} . This is the described system for which result HMR has really been proven. Although the model applies to the international MC market, the result is formulated as concerning the more specific described system with symmetric goods, homothetic preferences and so on. According to Krugman, 'The results were arrived at, however, only for a special case designed to make matters as simple as possible. Our next question must be the extent to which these results generalize' (1980, p. 958). Krugman knew that he was describing a 'missing system', but if the model could be generalised, then the specificities would not affect the result.

In a section entitled 'Generalizations and extensions' he relaxes the assumption that the population size is the same in the two countries (d_{size}) by letting the populations be of arbitrary size. He then shows that the HMR is unchanged (ibid.). In other words, this generalisation shows that a more general model M_K' yields the same result:

$$M_K' \vdash \text{HMR}(d_4(d_{mc}(d_{CES}, d_s, d_h, \dots))).$$

If the result crucially depended on the special assumption of equal population sizes, one would have reason to think that it only holds in circumstances that never hold in reality, namely that it is merely an artefact of tractability assumptions. By generalising the result in this way, Krugman showed that it is more likely to hold in reality. The main mechanisms, the target, the actual system and the main result remain the same, but they are described in terms of a more general set of assumptions. It has been shown that the functioning of the mechanism and its characteristic results are independent of some of the details included in the original description. The key change is that (c) the described system $d_4(d_{mc}(d_{CES}, d_s, d_h, d_{size}, \dots))$ in

the original model is a proper subset of the described system in the generalisation $d_4(d_{mc}(d_{CES}, d_s, d_h, \dots))$.⁵

M_K applies to S_2 because Krugman's target is $T = \{p_1, p_3, p_4\}$, the model represents those properties, and international monopolistic trade $S_2 = \{p_1, p_3, p_4, \dots\}$ has the properties defining the target. In other words, the generalised version of the proof of HMR does not generalise the model in terms of (a) applying it to new systems or phenomena. It generalises by *generalising the result* (c) because the new described system subsumes the old one, but not vice versa.

Figure 2 below depicts the various properties in and outside of the conceptualised system, and illustrates the different systems in Krugman's epistemically beneficial generalisation. This generalisation changes the described system from $d_4(d_{mc}(d_{CES}, d_s, d_h, d_{size}, \dots))$ to $d_4(d_{mc}(d_{CES}, d_s, d_h, \dots))$. The grey crescent-shaped area inside the conceptualised system describes the systems that are captured by the described system $d_4(d_{mc}(d_{CES}, d_s, d_h, \dots))$ but not by the described system $d_4(d_{mc}(d_{CES}, d_s, d_h, d_{size}, \dots))$, in other words, countries with different populations.

8.3 Another example: Wald's proofs

Let us consider the development of proofs of existence for general equilibria as described by Weintraub (1985) and Wald (1951). One such proof is given in Wald's 1935 paper under the restrictive assumption that all goods are independent (i.e., have zero cross elasticity of demand). Let x_i denote the demand for good i , and p_i its price: f is a monotonic function with $f' < 0$. Wald's initial assumption was:

$$x_i = f(p_i) \quad (W35)$$

Substitutes and complements do not belong to the set of independent goods. The first proof thus shows that there is an equilibrium, but the assumptions imply that goods are never substitutes for or complements of each other. The first proof thus effectively assumes that there are no markets with monopolistic competition nor markets with complementary goods. The second proof from 1936 removes the restriction that the existence of the equilibrium depends on there being no such interrelationships. Here Wald assumes:

$$x_i = f(p_1, p_2, \dots, p_i, \dots, p_n) \quad (W36)$$

Mäki (1992) notes that equation (W36) is more 'realistic' than equation (W35). The original model was already intended to apply to all markets. It is just that the 1935 model characterises some of those markets inaccurately. MC markets exist,

⁵ Lest this formulation leads the readers astray, generality is not a matter of the number of model descriptions. After all, as already discussed, Krugman's simpler formulation U is a special case of the more complicated DSP. I argue in the companion paper 'The epistemic benefits of generalisation in modelling II: Expressive power and abstraction', that epistemically beneficial assumptions typically decrease the number of assumptions in a model.

and Wald clearly wanted his 1935 result to apply to them even though mis-describing them.

Generalising the model shows that the assumption of independent goods, which was introduced for reasons of tractability, is not needed for deriving the existence of a general equilibrium. It was hoped that this assumption was not crucial for the result, and the 1936 generalisation confirms that it was not. Wald's original model clearly applied to general equilibria because it was able to characterise all markets simultaneously and provide a notion of equilibrium for them. The description of the market as consisting of independent goods is true of some markets, but this set is a small, proper subset of the set of all markets. Wald's original model thus applied to general equilibria, but the 1936 version generalised the result so that the described system was true about a much greater number of actual markets. Clearly, the generalised version generalises the result but it does not generalise the model in terms of (a) increasing the number of explained phenomena or increasing the number of systems to which the model applies specifically.

In both of these examples, generalising the result shows that the set of actual systems about which the model result is true is larger than before. Note carefully, however, that this set is different from the conceptualised system subsumed by the modeller's target.

9 Discussion and conclusions

I have discussed generalisation using examples from economics. A question that arises naturally is whether my results generalise to other model-heavy sciences. Tractability tends to compel economists to construct models that specify described systems they know to be merely special cases of what they really want to describe. There is thus a lot to generalise, and modellers know from the very beginning what would constitute a better model by comparing the conceptualised with the described system. In other words, if epistemically beneficial generalisations are particularly important in economics, it is because the problem of unrealistic assumptions is particularly important. At the same time, there is nothing in the account provided here that would limit its applicability to economics.

Distinguishing between conceptualised systems and targets allows for formulating the notion of generic model applicability which is able to recognise changes in the model descriptions as generalisations even when both models concern a given target and a given system.

Generic and specific applicability give rise to three kinds of generalisation: (a) applying a model specifically to new conceptualised systems (or increasing the number of phenomena the model explains), (b) abstracting the model, and (c) generalising results. Abstracting the model and applying a model to a new conceptualised system are not intrinsically epistemically beneficial in that they do not necessarily show that the model components or model results are more likely to be true. Instead, (a) and (b) have undesirable consequences. Abstraction means relaxing the standards of how precisely the target needs to be described, and as we have seen, applying

a model to a new system may decrease the generality of a model in terms of its generic applicability.

On the other hand, both (a) and (b) may be useful if they bring other benefits. In particular, if abstracting a model allows for proving a new important result about a larger set of conceptualised systems and if the new result concerning a new conceptualised system is important, then of course these new results may well be important enough to justify the instrumental usefulness of (a) or (b) for such results. Similar comments apply to case (a).

Generalising models bring epistemic benefits if they show that a model result also holds when some of the restrictive assumptions are removed. The epistemic benefit from a generalisation, when there is one, consists in showing that the model result does not depend on a particular idealisation because the generalisation removes it from the model. Insofar as there is uncertainty about whether the result is driven by such idealisations, the generalisation provides an epistemic boost to the model result. Epistemically beneficial generalisations can thus be attained when the model result remains the same.

Acknowledgements This paper started as a result of a set of events that transpired when Caterina Marchionni wrote a draft which pointed out how the epistemic benefits of generalisation are similar to those obtained by showing that a result is robust. I thought she was right, and I asked to join her in writing the paper. The paper was then re-written and its title was adopted while she was still a co-author. As time passed, however, we came to somewhat different views about how to develop the paper and I continued writing it alone. Due to the unusually central influence of one scholar in helping the writing process, I find it inappropriate to mention any others here, despite the fact that I can think of a few names. The anonymous reviewers have also been highly helpful. The usual disclaimer applies. The research was partly funded by a Nankai University Arts and Humanities Development Funds, grant no. ZB21BZ0218.

References

- Brakman, S., & Heijdra, B. J. (2004). Introduction. In S. Brakman & B. J. Heijdra (Eds.), *The monopolistic competition revolution in retrospect* (pp. 1–45). Cambridge University Press.
- Dixit, A. K., & Stiglitz, J. E. (1977). Monopolistic competition and optimum product diversity. *The American Economic Review*, 67(3), 297–308.
- Elliott-Graves, A. (2018). Generality and causal interdependence in ecology. *Philosophy of Science*, 85(5), 1102–1114.
- Enelow, J. M. (1981). Saving amendments, killer amendments, and an expected utility theory of sophisticated voting. *Journal of Politics*, 43(4), 1062–1089.
- Frigg, R. (2010). Fiction and scientific representation. In R. Frigg & M. C. Hunter (Eds.), *Beyond mimesis and convention* (pp. 97–138). Springer.
- Godfrey-Smith, P. (2006). The strategy of model-based science. *Biology and Philosophy*, 21, 725–740.
- Giere, R. N. (1988). *Explaining science: A cognitive approach*. University of Chicago Press.
- Hamminga, B. (1983). *Neoclassical theory structure and theory development*. Springer-Verlag.
- Jebeile, J., & Barberousse, A. (2016). Empirical agreement in model validation. *Studies in History and Philosophy of Science Part A*, 56, 168–174.
- Kitcher, P. (1989). Explanatory unification and the causal structure of the world. In P. Kitcher & W. C. Salmon (Eds.), *Scientific explanation, minnesota studies in the philosophy of science* (pp. 410–505). University of Minnesota Press.
- Knuuttila, T. (2017). Imagination extended and embedded: artifactual versus fictional accounts of models. *Synthese*. <https://doi.org/10.1007/s11229-017-1545-2>.

- Krugman, P. (1980). Scale economies, product differentiation, and the pattern of trade. *American Economic Review*, 70(5), 950–959.
- Kuorikoski, J., Lehtinen, A. & Marchionni, C. (2010). Economic modelling as robustness analysis. *British Journal for the Philosophy of Science*, 61(3), 541–567.
- Lehtinen, A. (2007). The welfare consequences of strategic voting in two commonly used parliamentary agendas. *Theory and Decision*, 63(1), 1–40.
- Lehtinen, A. (2015). Strategic voting and the degree of path-dependence. *Group Decision and Negotiation*, 24(1), 97–114.
- Levins, R. (1966). The strategy of model building in population biology. *American Scientist*, 54(4), 421–431.
- Levins, R. (1993). A response to Orzack and Sober: Formal analysis and the fluidity of science. *Quarterly Review of Biology*, 68(4), 547–555.
- Levy, A. (2015). Modeling without models. *Philosophical Studies*, 172(3), 781–798.
- Lewis, C. T., & Belanger, C. (2015). The generality of scientific models: A measure theoretic approach. *Synthese*, 192, 269–285.
- Mäki, Uskali. (1992). On the method of isolation in economics. In C. Dilworth (Ed.), *Intelligibility in science* (pp. 319–354). Atlanta and Amsterdam: Rodopi.
- Matthewson, J. (2011). Trade-offs in model-building: A more target-oriented approach. *Studies in History and Philosophy of Science Part A*, 42(2), 324–333.
- Matthewson, J., & Weisberg, M. (2009). The structure of tradeoffs in model building. *Synthese*, 170(1), 169–190.
- McKelvey, R. D., & Ordeshook, P. C. (1972). A general theory of the calculus of voting. In J. F. Herndon & J. L. Bernd (Eds.), *Mathematical applications in political science* (pp. 32–78). University Press of Virginia.
- Merrill, S. I. (1981). Strategic voting in multicandidate elections under uncertainty and under risk. In M. J. Holler (Ed.), *Power, Voting, and Voting Power* (pp. 179–187). Physica-Verlag.
- Neary, J. P. (2004). Monopolistic competition and international trade theory. In S. Brakman & B. J. Heijdra (Eds.), *The monopolistic competition revolution in retrospect* (pp. 159–184). Cambridge University Press.
- Odenbaugh, J. (2018). Models, models, models: A deflationary view. *Synthese*. <https://doi.org/10.1007/s11229-017-1665-8>.
- Ordeshook, P. C., & Palfrey, T. R. (1988). Agendas, strategic voting, and signaling with incomplete information. *American Journal of Political Science*, 32(2), 441–466.
- Orzack, S. H. (2012). The philosophy of modelling or does the philosophy of biology have any use? *Philosophical Transactions of the Royal Society B*, 367, 170–180.
- Orzack, S. H., & Sober, E. (1993). A critical assessment of Levins's the strategy of model building in population biology (1966). *Quarterly Review of Biology*, 68(4), 533–546.
- Parker, W. S. (2015). Getting (even more) serious about similarity. *Biology and Philosophy*, 30(2), 267–276.
- Potochnik, A. (2017). *Idealization and the aims of science*. Chicago University Press.
- Räz, T. (2017). The Volterra principle generalized. *Philosophy of Science*, 84(4), 737–760.
- Rol, M. (2008). Idealization, abstraction, and the policy relevance of economic theories. *Journal of Economic Methodology*, 15(1), 69–97.
- Stiglitz, J. E., & Dixit, A. K. (1993). Monopolistic competition and optimum product diversity: Reply. *American Economic Review*, 83(1), 302–304.
- Strevens, M. (2004). The causal and unification approaches to explanation unified—Causally. *Noûs*, 38(1), 154–176.
- Strevens, M. (2008). *Depth: An account of scientific explanation*. Harvard University Press, Cambridge, Mass.
- Thomson-Jones, M. (2010). Missing systems and the face value practice. *Synthese*, 172(2), 283–299.
- Wald, A. (1951). On some systems of equations of mathematical economics. *Econometrica*, 19(4), 368–403.
- Walliser, B. (1994). Three generalization processes for economic models. In B. Hamminga & N.B. De Marchi (Eds.), *Poznan Studies in the philosophy of the sciences and the humanities; idealization vi: idealization in economics* (pp. 55–69). Rodopi.
- Weintraub, E. R. (1985). *General equilibrium analysis: Studies in appraisal*. Cambridge University Press.

- Weisberg, M. (2004). Qualitative theory and chemical explanation. *Philosophy of Science*, 71(5), 1071–1081.
- Weisberg, M. (2007). Who is a Modeler? *British Journal for the Philosophy of Science*, 58(2), 207–233.
- Weisberg, M. (2012). Getting serious about similarity. *Philosophy of Science*, 79(5), 785–794.
- Weisberg, M. (2013). *Simulation and similarity: Using models to understand the world*. Oxford University Press.
- Weisberg, M. (2015). Biology and philosophy symposium on simulation and similarity: Using models to understand the world. *Biology and Philosophy*, 30(2), 299–310.
- Woodward, J., & Hitchcock, C. (2003). Explanatory generalizations, Part II: Plumbing explanatory depth. *Noûs*, 37(2), 181–199.

Publisher's Note Springer Nature remains neutral with regard to jurisdictional claims in published maps and institutional affiliations.