

TRANSMISSION FAILURE EXPLAINED*

MARTIN SMITH

University of Glasgow

In this paper I draw attention to a peculiar epistemic feature exhibited by certain deductively valid inferences. Certain deductively valid inferences are unable to enhance the *reliability* of one's belief that the conclusion is true – in a sense that will be fully explained. As I shall show, this feature is demonstrably present in certain philosophically significant inferences – such as GE Moore's notorious 'proof' of the existence of the external world. I suggest that this peculiar epistemic feature might be correlated with the much discussed phenomenon that Crispin Wright and Martin Davies have called 'transmission failure' – the apparent failure, on the part of some deductively valid inferences to *transmit* one's justification for believing the premises.

I. PRESERVATION AND TRANSMISSION

Fred Dretske briefly sent the epistemology world into a spin with his startling claim that knowledge is not always preserved by deduction (Dretske, 1970). Suppose that one knows that a proposition P is true, notices that proposition Q deductively follows from proposition P and concludes that proposition Q is true too. If one's belief that Q does not qualify as knowledge then we might say that the inference from P to Q fails to *preserve* the status of knowledge. Dretske's contention then is that knowledge is not preserved across all deductively valid inferences. This is more commonly expressed as the claim that knowledge is not 'closed under' deductive inference. I shall stick with the term 'preservation'. Deductive inferences are sometimes defined

* Thanks go to Juan Comesaña, Martin Davies, Andy Egan, James Pryor, Peter Roeper, Ernest Sosa and Crispin Wright for valuable comments upon earlier incarnations of this paper. I am also indebted to audiences at the University of Western Australia, the University of York and the University of St Andrews.

as those that always *preserve* truth. My use of the term is essentially continuous with this usage. Dretske's contention could be expressed in this way: Some inferences that always preserve the status of truth do not always preserve the status of knowledge.

Dretske supported his contention using the following now famous example: Suppose I am at the zoo and spy some zebras grazing in a nearby enclosure. Noting their black and white, striped, equine appearance I proceed to reason as follows:

- (1) Those animals are zebras.
- (2) If those animals are zebras then they are not mules disguised to look like zebras.
- (3) Therefore, those animals are not mules disguised to look like zebras.

According to Dretske, I can know that the first premise is true on the basis of the visual evidence at my disposal. Presumably some low-level background information about zebras and mules is sufficient for me to know the second premise. The inference is palpably valid. However, when it comes to the conclusion, my evidence for accepting the first premise is rendered impotent or inert – it fails to discriminate between the conclusion and its most salient alternatives. According to Dretske, if I believe that the animals before me are not disguised mules on the basis of their black and white, striped, equine appearance, then this belief will not qualify as knowledge.

Dretske's construal of this case remains highly controversial however and, aside from a few notable exceptions (Nozick, 1981, pp204-211, Heller, 1999), I think it is fair to say that Dretske's contention has not been embraced by the wider epistemological community. For most epistemologists, the idea that deductive inference represents an epistemically safe way of extending one's belief corpus is

simply too dear. And yet, it seems that Dretske has certainly put his finger on *something*. There is something undeniably suspicious about the zebra inference.

More recently, Crispin Wright (1985, 2000, 2002, 2003, 2004) and Martin Davies (1998, 2000, 2003a, 2003b, 2004) have drawn attention to a distinct, though closely related, epistemic misadventure they term *transmission-failure*. Suppose that one *justifiably believes* that a proposition P is true, notices that proposition Q deductively follows from proposition P and concludes that proposition Q is true too. If one's belief that Q fails to qualify as justified *in virtue of the inference from P* then, we might say, the inference fails to *transmit* the status of justification. One's belief that Q may still be justified, of course – but it must be justified in virtue of something other than the inference from P. As it is sometimes put, a valid inference fails to transmit the status of justification iff I cannot *earn* or procure a justification for believing the conclusion by deploying my justification for the premise and drawing the inference.

Both Davies and Wright propose that inferences such as Dretske's zebra inference ought to be regarded not as cases of *preservation* failure, but rather as cases of *transmission* failure. It is implausible to think that my belief that the animals are not disguised mules could be justified on the basis of my visual evidence for believing that they are zebras. But it doesn't follow from this that I have no justification *whatsoever* for believing that the animals are not disguised mules. Perhaps my justification issues from some *other* source. Perhaps it is supplied by low-level collateral information about, say, the trustworthiness of zoos. Perhaps it is simply in place by default.

An inference can, of course, transmit certain kinds of justification while failing to transmit others (see Wright, 2000, pp141). If, for instance, I believed that the animals before me are zebras on the basis of a DNA test, then *this* sort of justification would presumably flow quite unimpeded through the zebra inference. It's only when I believe that the animals are zebras on the basis of indiscriminate visual information that a problem seems to arise.

Here is another example:

- (1) I counted nine dollars on the table.
- (2) There are nine dollars on the table.
- (3) Therefore, I did not miscount the money on the table.

If I believe that there are nine dollars on the table on the strength of my *counting* nine dollars on the table – that is, if I believe the second premise on the strength of the first premise – then this reasoning looks decidedly suspect and, plausibly, represents a case of transmission failure. If, on the other hand, I believe that there are nine dollars on the table on the basis of, say, a friend's testimony, then my reasoning may be in perfectly good order.

Transmission failure is an idea that divides epistemologists. Some – notably Silins (2005) – argue that there are no genuine examples at all. Further, even amongst those who accept that there are genuine examples of transmission failure, there is relatively little agreement as to just what these examples are. While Dretske's zebra inference is widely regarded as a plausible case, the transmitting status of certain other, especially philosophical, inferences is more contentious. One inference that has

sparked particular controversy is G.E. Moore's notorious 'proof' of the existence of the external world (see Moore, 1939).

Moore, the indignant advocate of so-called 'common sense' epistemology, held his hands before his face and declared:

- (1) I have hands.
- (2) If I have hands then the external world exists.
- (3) Therefore, the external world exists.

When first exposed to Moore's reasoning most find it exceedingly dissatisfying. As it turns out, though, it is not all that easy to pinpoint exactly where Moore goes wrong. As Moore himself stresses, the inference is palpably valid and it looks as though one does have justification for endorsing the premises. Most people, indeed, would be more confident of Moore's premises than of any premise to which an external world sceptic might appeal. It's not obvious what more one should demand.

Both Wright and Davies propose that the problem is transmission failure – one could not be justified in believing in the existence of the external world in virtue of drawing the inference Moore suggests (Wright, 1985, pp434-438, 2002, pp336-367, Davies, 1998, pp350, 2004). Tyler Burge (2003) and James Pryor (2004), however, take exception to this diagnosis. According to Burge and Pryor, the problem with Moore's argument is *dialectical* rather than epistemic. The argument may be *unpersuasive*, but it could, nevertheless, supply one with a justification for accepting its conclusion.

In this paper I identify an unusual feature that a deductively valid inference might exhibit – namely, an inability to enhance the reliability of one’s belief that the conclusion is true. This is a feature that is demonstrably exhibited by some valid inferences and, indeed, demonstrably exhibited by Moore’s inference. I argue that this feature is sufficient for transmission failure, as characterised by Wright and Davies. If I am right, then not only will the views expressed by Silins and by Burge and Pryor be refuted but we will have at our disposal a plausible explanation as to how transmission failure ‘works’ – as to the principles that underlie it.

Before proceeding, it is worth making the following observation: Although preservation failure tends to be discussed in relation to knowledge and transmission failure tends to be discussed in relation to justification, this is something of an historical accident. Suppose one’s belief that a proposition *P* is true has a certain status *S*. Suppose one notices that proposition *Q* deductively follows from proposition *P* and concludes that proposition *Q* is true. If one’s belief that *Q* fails to qualify for status *S*, we might say that the inference fails to *preserve* *S*. If one’s belief that *Q* fails to qualify for status *S* in virtue of the inference from *P*, we might say that the inference fails to *transmit* *S*. Both knowledge and justification, and indeed a number of other things besides, can be meaningfully substituted for *S* in both the preservation and transmission schemas.

What we have then, are four properties of interest – namely, knowledge preservation, justification transmission, knowledge transmission and justification preservation – of which the first two have been most widely discussed in the literature. All four, however, will feature in the remainder of this paper.

II. SAFETY

The next character in this story is Ernest Sosa. In a series of influential papers (1999a, 1999b, 2000, 2002) Sosa offers a particular construal of a rather intuitive idea – namely, in order for one to know that a proposition *P* is true, it is necessary that one's belief that *P* be held on a *safe* or secure basis. In Sosa's view, what it means for one proposition – say, (B) *the animals are black and white, striped and equine* – to serve as a safe basis upon which to believe another proposition – say, (P) *the animals are zebras* – is for the two propositions to be distributed in the right sort of way throughout certain important possible worlds. B provides a safe basis upon which to judge that P just in case, in all those possible worlds in which B is true and which are *very similar to* or closely resemble the actual world, P is true too¹. In other words, one has to depart gratuitously from actuality in order to find possible worlds in which B is true and P is false. A belief held upon a safe basis might be described as a *safe* belief.

It is worth pointing out that, provided that B is true, the actual world will itself be one of the very similar worlds in which B is true. The actual world is, presumably, very similar to itself. It follows from this that safe beliefs have to be true beliefs. If B

¹ Sosa claims that B provides a safe basis upon which to believe that P just in case the following subjunctive conditional holds true: 'B would not be true without P being true' (see, for instance, Sosa, 1999a, section E) or 'Provided that B is true, P would be true too'. The above possible worlds condition is endorsed by Sosa as being a plausible semantic analysis of these subjunctive conditionals. I have opted to simply do away with the middleman here. The possible worlds analysis of safety is perfectly intuitive in its own right (at least for those conversant in the parlance of possible worlds), and moving directly to it allows us to avoid potentially distracting issues about the semantics of subjunctive conditionals.

is a safe basis upon which to believe that P and one believes that P on basis B, then P must hold true.

Sosa claims that, in order for one to know that a proposition P is true, it is necessary that one's belief that P be held upon a safe basis, in the manner just defined. This claim has been endorsed by others – such as Williamson (2000, section 5.3) and Pritchard (2002, 2005). Still others defend theories of knowledge that appear to have Sosa's claim – or something quite like it – as a consequence (see for instance Dretske, 1971, Goldman, 1976, Nozick, 1981, chap. 3, DeRose, 1995). I offer no arguments in support of Sosa's claim – unless, of course, its power to explain the phenomenon of transmission failure be thought to weigh in its favour, in which case this entire section could be seen as a rather elliptical argument of this kind.

Suppose, once again, that I am at the zoo, spy some zebras grazing in a nearby enclosure, take note of their black and white, striped, equine appearance and form the belief that they are zebras. Clearly, there will be some very similar possible worlds in which the animals before me are not zebras – worlds in which they are, say, pigs or lions or macaques. In none of these worlds, however, will the animals exhibit a black and white, striped equine appearance. Of course, there will be possible worlds in which the animals appear to be black and white, striped and equine without being zebras – worlds in which the animals are disguised mules for instance. Provided, though, that I'm visiting a regular, responsible zoo with effective security measures etc. these worlds will be comparatively remote from actuality, in which case my belief will qualify as safe.

If B serves as a safe basis upon which to believe that P then, given the above characterisation of safety, B must also serve as a safe basis upon which to believe any deductive consequence of P. A deductively inferred belief cannot be in any greater danger of being false than the belief from which it was deductively inferred.

Suppose that, having judged that (P) the animals are zebras, on the basis that (B) the animals appear to be black and white, striped and equine, I proceed to infer that (Q) the animals are not disguised mules. As we have already seen, the most similar worlds in which the animals are disguised mules – that is, the most similar worlds in which Q is false – are too dissimilar to be taken into consideration for the purposes of evaluating safety. It follows immediately that B is a safe basis upon which to judge that Q. There won't be any sufficiently similar worlds in which B is true and Q is false simply because there won't be any sufficiently similar worlds *in which Q is false*.

One might be tempted by the idea that our standards of comparative similarity are not set in stone but, rather, vary in coarseness from context to context taking our standards of safety along with them. Maybe so. But this much, I suggest, is clearly true: Provided we hold our standards of safety constant, deductive inference can never take us from a safe belief to an unsafe belief. Suppose one safely believes that P is true, notices that Q is a deductive consequence of P and proceeds to infer that Q is true too. If one's belief that P is safe, then one's belief that Q must be safe. Safety, we might say, is *preserved* by deductive inference. If Dretske is right and *knowledge* is not always preserved, then this will not be due to the safety condition.

As the zebra example helps to make clear, it is possible to distinguish two rather different ways in which a belief might qualify as safe – that is, two different ways in which a proposition B might count as a safe basis upon which to believe a further proposition P. First, B might count as a safe basis upon which to believe that P in virtue of the modal relationship between B and P. This is the case when B is *the animals are black and white, striped and equine* and P is *the animals are zebras*. Second, B might count as a safe basis upon which to believe that P purely in virtue of the modal profile of P. This is the case when B is *the animals are black and white, striped and equine* and P is *the animals are not disguised mules*. Say, in the former case, that B is a *contributing* safe basis upon which to believe that P and, in the latter case, that B is a *non-contributing* safe basis upon which to believe that P. A belief held upon a contributing safe basis might be described as *safe in virtue of its basis*, while a belief held upon a non-contributing safe basis might be described as *safe purely in virtue of its content*.

B is a safe basis upon which to believe that P just in case in all the most similar worlds in which B is true, P is true too. B is a contributing safe basis upon which to believe that P just in case, in addition, in all the most similar worlds in which P is *false*, B is false too. B is a non-contributing safe basis upon which to believe that P just in case, in some of the most similar worlds in which P is false, B is still true. See figs. 1 and 2 below:

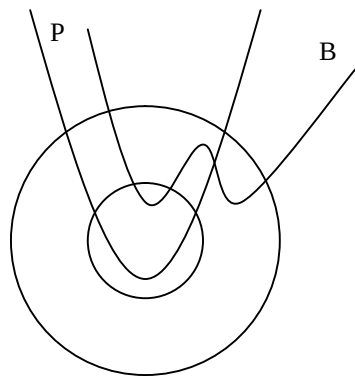


Fig. 1

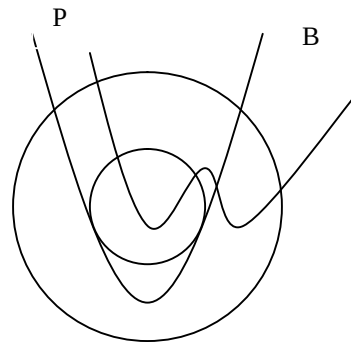


Fig. 2

The inner circle in each diagram represents the class of possible worlds that are very similar to the actual world, while the outer ring represents the class of possible worlds that rank next in terms of similarity to the actual world. Figure 1 represents a situation in which B qualifies as a contributing safe basis upon which to believe that P, while figure 2 represents a situation in which B qualifies as a non-contributing safe basis upon which to believe that P. Notice that the distinction between contributing and non-contributing safe bases, as defined here, effectively breaks down when it comes to necessary truths. As a kind of limiting case, any basis will qualify, according to the characterisation given above, as a contributing safe basis upon which to believe a necessary truth.

If B is a contributing safe basis upon which to believe that P, it does not follow that B is a contributing safe basis upon which to believe a deductive consequence Q of P. A belief may be safe purely in virtue of its content even though it was deductively inferred from a belief that is safe in virtue of its basis. The zebra inference, of course, provides one example of this phenomenon. The situation is represented in fig 3:

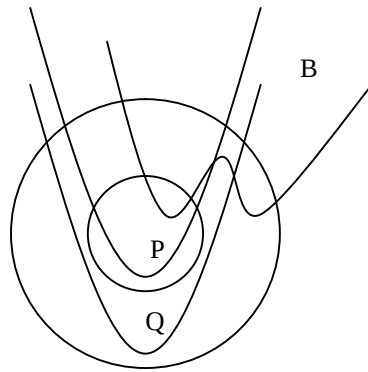


Fig. 3

Suppose one safely believes that P is true, notices that Q is a deductive consequence of P and proceeds to infer that Q is true too. If one's belief that Q is safe purely in virtue of its content, then presumably it cannot be said to be safe in virtue of the inference from P . One's belief that Q would have been safe irrespective of whether one inferred Q from P – it would have been safe even if it were held as an article of faith. Safety, we might say, is not always *transmitted* by deductive inference². When knowledge is not transmitted, this may be due to the safety condition.

Suppose one knows that P is true, notices that Q is a deductive consequence of P and proceeds to infer that Q is true too. Suppose, further, that one's belief that Q is safe purely in virtue of its content. Given that safety is one constituent of knowledge, one's belief that Q could not be said to qualify as knowledge in virtue of the inference. In other words, the inference does not transmit knowledge from premise to

² If we formulate the safety condition as a subjunctive conditional – provided B is true, P would be true too ($B \Box \rightarrow P$) – then the preservation of safety across deductive inference is ensured by the fact that *weakening the consequent* is a valid inference pattern for such conditionals: $B \Box \rightarrow P$ and $P \rightarrow Q$ entail $B \Box \rightarrow Q$ (' $\rightarrow$ ' should be read here as a strict conditional). The *transmission* of safety, however, can fail precisely because *strengthening the antecedent* is not a valid inference pattern for subjunctive conditionals: $\sim Q \rightarrow \sim P$ and $\sim P \Box \rightarrow \sim B$ do not entail $\sim Q \Box \rightarrow \sim B$ (see Lewis, 1979, section 1.8).

conclusion. It should be emphasised that safety is only a *necessary* condition for knowledge and that, even though one's belief that Q is safe purely in virtue of its content, the inference from P may yet play some role in ensuring that further conditions for knowledge are satisfied. I don't wish to rule out such a possibility. However, even if the inference is playing some such role, it would be wrong to claim that one's belief that Q qualifies as knowledge *in virtue of it*.

It may be necessary to say a little more about the 'in virtue of' relation invoked here. The relation that I have in mind is of an explanatory kind. When I say that something has property F in virtue of having property G, what I mean, roughly speaking, is that the possession of G serves to *explain* the possession of F. Suppose that one knows that Q. I take it that the safety of one's belief that Q is part of what explains its qualifying as knowledge. If the safety of one's belief that Q is fully explained by Q's modal profile, then this, in turn, will form part of the explanation as to why one's belief that Q qualifies as knowledge. The knowledge status of one's belief could not then be explained *exclusively* in terms of the inference that one has drawn and the resultant basis of one's belief.

All that I am claiming here is that a failure to transmit safety is a sufficient condition for a failure to transmit knowledge. Even if an inference transmits safety, it may yet fail to transmit knowledge because of a failure to transmit some other necessary precondition. For present purposes I will leave this possibility open.

It is possible to make *comparative* as well as absolute judgments about the safety of beliefs and bases. That is, it makes perfect sense to say things like: 'My

belief that the wall *looks* red is safer than my belief that the wall *is* red' or 'It is safer to believe that the culprit had blonde hair on the basis of twelve people's testimony than on the basis of one person's testimony'.

Within the framework I have developed there is a perfectly natural way to make sense of such comparisons. A belief that P is safer than a belief that Q just in case the most similar worlds in which Q is false and its actual basis true are *more similar* to the actual world than the most similar worlds in which P is false and its actual basis true. B is a safer basis upon which to believe that P than C is just in case the most similar worlds in which C is true and P is false are more similar to the actual world than the most similar worlds in which B is true and P is false. B might be described as a safer basis than C upon which to believe that P just in case we have to depart further from actuality in order to find worlds in which B is true and P is false than we do to in order find worlds in which C is true and P is false. It is important to note that B and C could still both qualify as *safe bases*. That is, it could still be the case that the most similar worlds in which B is true and the most similar worlds in which C is true are worlds in which P is true.

This notion of comparative belief safety effectively coincides with Keith DeRose's notion of the *comparative strength* of one's epistemic position (see DeRose, 1995). According to DeRose, one is in a stronger epistemic position with respect to a proposition P than a proposition Q just in case those possible worlds in which one believes that P falsely via one's actual method are less similar to the actual world than those worlds in which one believes that Q falsely via one's actual method (see DeRose, 1995, section 12). DeRose exploits this notion in developing a certain

contextualist account of knowledge attributions. I'll have a little more to say about DeRose's theory in the next section.

Certain beliefs – such as the belief that the animals before me are not disguised mules – enjoy a certain *default* level of safety in virtue of their content alone. Imagine once again the zoo scenario. By basing my belief that (Q) the animals before me are not disguised mules, upon the fact that (B) they appear to be black and white, striped and equine, it is clear that I fail to make my belief any *safer*. That is, I fail to enhance the safety of my belief beyond the default level that it enjoys in virtue of its content. The most similar worlds in which B is true and Q is false are simply amongst the most similar worlds in which Q is false. My belief that Q would have been just as safe even if it had been held as an article of faith.

Now suppose I believe that (Q) the animals before me are not disguised mules on the basis that (C) a DNA test yielded the result that they are zebras. The most similar worlds in which C is true and Q is false will be worlds at which not only are the animals disguised mules, but the DNA test has been bungled or the results have been tampered with or some such. Provided the DNA test is actually conducted in a controlled, rigorous manner, these possible worlds will be less similar to the actual world than the most similar worlds in which the animals are disguised mules.

Although my belief that (Q) the animals before me are not disguised mules, enjoys a certain level of safety in virtue of its content, by basing my belief upon the fact that (C) a DNA test yielded the result that they are zebras, I *enhance* the safety of my belief beyond this default level. The safety of my belief that Q is *overdetermined*

as it were. It is safe *both* in virtue of its basis *and* in virtue of its content. C is a safe basis upon which to believe that Q both in virtue of the modal profile of Q and in virtue of the modal relationship between C and Q.

Consider again an earlier example: Suppose I carefully count the money on the table before me and arrive at the correct result of nine dollars. Suppose I believe, on the basis of my counting that (P) there are nine dollars on the table, and that (Q) I counted nine dollars on the table. Suppose I then infer, from P and Q, that (R) I did not miscount the money on the table. My belief that R is true is plausibly imbued with a certain default safety in virtue of its content alone. However, by inferring R from P and Q and, thus, believing it on the basis of my counting, I fail to enhance the safety of my belief beyond the default level. The most similar worlds at which I believe falsely that I didn't miscount on the basis of my counting are simply the most similar worlds at which I miscount.

Now suppose that both I and my friend carefully count the money and arrive at the correct result of nine dollars. If I believe that (P) there are nine dollars on the table, on the basis of my counting *and* my friend's testimony and infer R as before, then I *do* manage to enhance the safety of my belief beyond its default level. The most similar worlds in which my basis leads me astray will be worlds in which I miscount *and* my friend miscounts (or perhaps lies). Provided my friend is honest and counts carefully, these worlds will be less similar to the actual world than the most similar worlds in which I miscount.

I claimed above that B is a non-contributing safe basis upon which to believe a proposition P, just in case B is a safe basis upon which to believe that P *purely* in virtue of the modal profile of P. The ‘purely’ turns out to be important. If B is a safe basis upon which to believe that P *both* in virtue of the modal profile of P *and* in virtue of the modal relationship between B and P, then B will qualify as a *contributing* safe basis upon which to believe that P, in line with the definition I offered. Such a situation is represented in fig. 4:

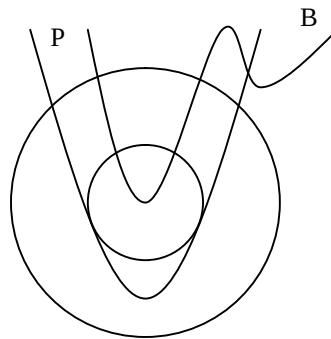


Fig. 4

I have offered the following account of the transmission failure of knowledge: Suppose that one knows that P, notices that Q deductively follows from P and concludes that Q is true too. If one’s belief that Q turns out to be safe purely in virtue of its content, then this inference fails to transmit the status of knowledge. Naturally, the ‘purely’ qualification is crucial here too. If one’s belief that Q is safe both in virtue of its content *and* in virtue of its basis then knowledge should be deemed to be transmitted.

This account of transmission failure is equivalent to the following: Suppose that one knows that P, notices that Q deductively follows from P and concludes that Q

is true too. If, by drawing the inference, one fails to *enhance* the safety of one's belief that Q beyond whatever default level it enjoys in virtue of its content, then this inference fails to transmit the status of knowledge. There is one proviso – namely, that it is *possible* to enhance the safety of one's belief that Q or, alternately, that there are possible worlds in which Q is false. If Q is a necessary truth, then a failure, on the part of an inference, to enhance the safety of one's belief that Q need not be indicative of a failure to transmit knowledge.

When it comes to a belief in the existence of the external world, it is clear that we are dealing with an extraordinarily potent default level of safety – far more potent, presumably, than a belief to the effect that the animals before me are not disguised mules. There may be possible worlds in which reality is mind-dependent in some relevant sense but, provided the actual world is not amongst them, they should, surely, be regarded as very dissimilar. Clearly, if I were to base my belief that the external world exists upon the appearance of hands before my face, this would not in the least enhance the safety of my belief beyond this default level. In those possible worlds in which reality is mind dependent in the relevant sense, perceptual experiences will not have external causes. Rather, they will be explicable in some other way – perhaps as a spontaneous expression of one's mind. In any case, the most similar worlds in which reality is mind dependent *and* there appear to be hands before my face will simply be amongst the most similar worlds in which reality is mind dependent. My account predicts, then, that Moore's inference fails to transmit knowledge from premises to conclusion.

It may be that there is literally nothing that we can do in order to enhance the safety of our commitment to the existence of the external world. That is, it may be that we could never identify a contributing safe basis upon which to believe in the existence of the external world. If so, I suggest that this is just as much a reflection upon the belief's extraordinary default level of safety as it is upon our own epistemic shortcomings, in which case it is not clear that this should be regarded as a source of any epistemic embarrassment.

III. RELIABILITY

A belief need not be safe in order to be justified. This follows straightforwardly from the observation that safe beliefs have to be true while justified beliefs do not. Suppose that the animals before are, in actual fact, disguised mules, but I take them to be zebras on the basis of their black and white, striped, equine appearance. In this case, my belief is clearly unsafe but, provided I am not aware of any evidence suggesting the possibility of deception, it will still be perfectly justified.

I believe, however, that there is a condition upon justified belief that is, in important respects, *analogous* to the safety condition upon knowledge. It is this: In order for a belief to qualify as justified it is necessary that it be held on a *reliable* or dependable basis. In my view, what it means for one proposition – say, (B) *the animals are black and white, striped and equine* – to serve as a reliable basis upon which to believe another proposition – say, (P) *the animals are zebras* – is for the two propositions to be distributed in the right sort of way throughout certain important possible worlds. However, rather than selecting these worlds on the basis of their

similarity to the actual world, as we might if we were evaluating the safety of B as a basis for believing P, we would be better served by selecting these worlds according to their *normalcy*, from the perspective of the actual world. B provides a reliable basis upon which to judge that P just in case, in all those possible worlds in which B is true and which are as normal, from the perspective of the actual world, as the truth of B permits them to be, P is true too³. In other words, one has to depart gratuitously from ideal normalcy in order to find possible worlds in which B is true and P is false. A belief held upon a reliable basis might be described as a *reliable* belief.

One might, of course, find the idea of comparative world normalcy rather obscure or extravagant. But, much like the idea of a possible world itself, the idea of comparative world normalcy sounds more extravagant than it is. I would say that anyone who has ever worked with a somewhat simplified or idealised approximation of a potentially complex actual situation has already encountered the idea of a (miniature) normal world, much as I intend it. For present purposes, I will not say anything more substantial about world normalcy. Indeed, I regard it as a virtue of my account of reliability that it is compatible with a range of different views about this. For further discussion of the relevant notion of normalcy see Smith (forthcoming).

It is worth pointing out that, even if basis B is true, the actual world need *not* be amongst the most normal worlds at which B is true. On any reasonable construal of normalcy, plenty of abnormal circumstances do prevail. It follows from this that

³ The above possible worlds condition might be endorsed as a plausible semantic analysis of the following *normic* conditional: Provided B is true, P would normally be true too. If this analysis is correct, then we can claim that B serves as a reliable basis upon which to believe that P iff such a normic conditional holds true. I won't discuss the semantics of normic conditionals here – but see Smith (2007).

reliable beliefs need not be true. If a belief is reliable then it must have a particular *disposition* – namely, a disposition to be true under normal conditions or, more strictly, under the most normal conditions in which it is held upon its actual basis. A reliable belief can perfectly well be false, provided that prevailing conditions are abnormal or aberrant.

I claim that, in order to be justified, a belief must be held upon a reliable basis in the sense just defined. While many epistemologists would endorse the sentiment that a justified belief must be held upon a reliable basis (see, for instance, Goldman, 1979, 1986, Alston, 1988, Schmitt, 1992, Chase, 2004), my construal of this condition is somewhat distinctive – though perfectly in keeping with the ordinary use of the term ‘reliability’. As with Sosa’s claim regarding the safety condition upon knowledge, I won’t offer any considerations buttressing my claim regarding the reliability condition upon justified belief – beyond pointing out that it has a certain intuitive appeal and affords a satisfying explanation of a phenomenon that many acknowledge – namely transmission failure.

Suppose that I am at the zoo, spy some animals grazing in a nearby enclosure, take note of their black and white, striped, equine appearance and form the belief that they are zebras. Suppose that the animals I see are, in actual fact, mules painted to look like zebras. Clearly, there will be some very similar possible worlds in which the animals before me are not zebras even though they exhibit a black and white, striped equine appearance – indeed, the actual world is one of these. Provided, though, that mule disguising is a comparatively unusual or abnormal practice at the actual world, my belief will still qualify as *reliable*. At none of the most *normal*

worlds in which the animals before me exhibit a black and white, striped, equine appearance will they be painted mules – one incident of mule painting is not going to change that. A black and white, striped, equine appearance will still be a reliable basis upon which to believe that an animal is a zebra, even though it happens to lead me astray under the unusual prevailing circumstances.

If mule disguising were not abnormal or unusual – if, say, it were an entrenched and routine practice – then my belief would not qualify as reliable. At such a world, a black and white, striped, equine appearance would not qualify as a reliable basis upon which to classify an animal as a zebra. Or imagine, alternately, a possible world in which some other kind of animal exhibited a black and white, striped equine appearance and could only be distinguished from a zebra in virtue of further, more subtle, features. In a world such as this I could not reliably classify an animal as a zebra on the basis of a black and white, striped, equine appearance, and would not be justified in doing so. If, in the actual world, I were to classify an animal as a magpie solely on the basis of its small, black and white, avian appearance, my belief would not qualify as reliable – since plenty of other birds, such as magpie larks, meet this description.

If B serves as a reliable basis upon which to believe that P then, given my characterisation of reliability, B must also serve as a reliable basis upon which to believe any deductive consequence of P. It cannot be any less appropriate for me to rely upon the truth of a deductively inferred belief than it is for me to rely upon the truth of the belief from which it was inferred.

Suppose that, having judged that (P) the animals are zebras, on the basis that (B) the animals appear to be black and white, striped and equine, I proceed to infer that (Q) the animals are not disguised mules. As we have already seen, the most normal worlds in which the animals are disguised mules – that is, the most normal worlds in which Q is false – are less normal than the most normal worlds in which B is true and, thus, too abnormal to be taken into consideration for the purposes of evaluating reliability. It follows immediately that B is a reliable basis upon which to judge that Q. There won't be any sufficiently normal worlds in which B is true and Q is false simply because there won't be any sufficiently normal worlds *in which Q is false*.

This remains true even if Q is false at the actual world and I am, in fact, looking at disguised mules. Even if the animals before me are disguised mules, it is still quite appropriate for me to rely upon the supposition that they are not. This supposition would be true in all of the most normal worlds in which the animals appear as they do.

One might think that our standards of comparative normalcy are inclined to vary in coarseness from context to context and that our standards of reliability may shift accordingly. Provided, though, we hold our standards of reliability constant, deductive inference can never take us from a reliable belief to an unreliable belief. Suppose one reliably believes that P is true, notices that Q is a deductive consequence of P and proceeds to infer that Q is true too. If one's belief that P is reliable then one's belief that Q must be reliable. Reliability, like safety, is *preserved* by deductive inference.

Reliability, of course, has another thing in common with safety – namely, it can originate both in the modal relationship between a belief and its basis and in the intrinsic modal profile of a belief's content. A proposition B might count as a reliable basis upon which to believe that P in virtue of the modal relationship between B and P. This is the case if B is the proposition *the animals appear to be black and white, striped and equine* and P is the proposition *the animals are zebras*. The class of ideally normal worlds will include worlds in which B is true, worlds in which B is false, worlds in which P is true and worlds in which P is false. However, the truth of B, in combination with the falsity of P, will implicate a departure from ideal normalcy.

Alternately, a proposition B might count as a reliable basis upon which to believe that P in virtue of the modal profile of P. This is the case if B is the proposition *the animals appear to be black and white, striped and equine* and P is the proposition *the animals are not disguised mules*. In this case, the truth of P is partially constitutive of what it is for conditions to be 'ideally normal'. The falsity of P *tout court* implicates a departure from ideal normalcy.

B is a reliable basis upon which to believe that P just in case in all the most normal worlds in which B is true, P is true too. B is a contributing reliable basis upon which to believe that P just in case, in addition, in all the most normal worlds in which P is *false*, B is false too. B is a non-contributing reliable basis upon which to believe that P just in case, in some of the most normal worlds in which P is false, B is still true. Suppose we re-interpret figs. 1 and 2 above such that the inner circle

represents the class of worlds that are ideally normal from the perspective of the actual world and the outer ring represents the class of worlds that rank next with respect to comparative normalcy. Reinterpreted thus, figs. 1 and 2 represent, respectively, a situation in which B qualifies as a contributing reliable basis upon which to believe that P and a situation in which B qualifies as a non-contributing reliable basis upon which to believe that P. The distinction between contributing and non-contributing reliable bases also breaks down when it comes to necessary truths. If P is a necessary truth then any basis will qualify as a contributing reliable basis, as defined here, upon which to believe that P.

If B is a contributing reliable basis upon which to believe that P, it does not follow that B is a contributing reliable basis upon which to believe every deductive consequence of P. A belief may be reliable purely in virtue of its content even though it was deductively inferred from a belief that is reliable in virtue of its basis. Reliability, then, is not always transmitted by deductive inference⁴. If a deductive inference fails to transmit reliability, then this might account for a failure to transmit justification.

All that I am claiming here is that a failure to transmit reliability is a sufficient condition for a failure to transmit justification. Even if an inference transmits reliability, it may yet fail to transmit justification because of a failure to transmit some other necessary precondition. For present purposes I will leave this possibility open.

⁴ If we formulate the reliability condition as a normic conditional – provided B is true, P would normally be true too ($B \blacksquare \rightarrow P$) – then the guaranteed preservation of reliability, along with the possibility of its non-transmission, can be explained by citing certain logical features of this conditional. In particular, they can be explained by the fact that weakening the consequent is, while strengthening the antecedent is not, a valid pattern. The logics governing subjunctive and normic conditionals do diverge, but they share this much (see Smith, 2007).

Crispin Wright (2000, 2004), Martin Davies (2000, 2004) and Fred Dretske (2000) all make use of a notion they term ‘negative entitlement’. Though they differ over the details, a negative entitlement is intended, in essence, to be a kind of epistemic credential that does not reflect any cognitive accomplishment on the part of a subject – either empirical or a priori. A negative entitlement is meant to be a sort of justification that one need not *earn* or procure. I believe that this notion – or something close to it at any rate – can be constructed within the framework I have developed.

We might say that one’s belief is *negatively justified* just in case (i) the belief is justified and (ii) the reliability condition upon justification is taken care of exclusively by the belief’s content. More precisely, a belief might be described as negatively justified just in case it is justified and held on a non-contributing reliable basis. If this condition is met then, clearly, one deserves no credit for the *reliability* of one’s belief. Since reliability is only a necessary condition for justification, it does not follow that one deserves no credit for the *justificatory status* of one’s belief – one might still deserve credit for ensuring that further necessary conditions are satisfied. We might say, more cautiously, that if one’s belief is negatively justified, in the sense just defined, then one cannot take *full* credit for its justificatory status.

When one infers a deductive consequence of a justified belief, there is no danger that the resultant belief will be unjustified. Some deductively valid inferences, however, can cause the *valence* of one’s justification to flip – to switch from positive to negative. When one infers a deductive consequence of a justified belief, there may

be a danger that the justificatory status of the resultant belief is something for which one cannot legitimately claim full credit. If one has a positive justification for believing a proposition P, notices that Q is a deductive consequence of P and comes to believe that Q is true too, then this inference fails to *transmit* reliability just in case it fails to *preserve* the status of positive justification.

If one holds a negatively justified belief that P, one may not be in a position to offer any considerations or reasons in favour of P – to defend the belief against those who might challenge one’s right to hold it. Citing a non-contributing reliable basis for believing P is not a way of defending the belief – it cannot make P seem any more plausible than it seems already. Asserting ‘The animals are not mules disguised to look like zebras – after all, they appear to be black and white, striped and equine’ is no more rationally persuasive than simply insisting ‘The animals are not mules disguised to look like zebras’ (perhaps it is *less* rationally persuasive). In contrast, if one holds a positively justified belief that P, then one should, in principle, have resources upon which to draw in defending one’s belief. Citing a contributing reliable basis for believing P can make P seem more plausible, at least for an audience with the right sorts of background beliefs.

One might, of course, find the very idea of negative justification to be counterintuitive. After all, it seems intuitively improper to claim justification for believing P unless one is in a position to offer considerations or reasons that support it. I suspect that this intuition is not quite as robust as is often supposed – but I won’t pursue this here. For now I wish to emphasise that the intuitive costs of *rejecting* the idea negative justification, in the precise sense defined here, are themselves

substantial. If a deductive inference fails to transmit reliability, then it fails to preserve positive justification. If there are genuine cases in which reliability is not transmitted by deductive inference, then we are confronted with a dilemma: Either (i) there is such a thing as negative justification or (ii) there are genuine cases in which justification *simpliciter* is not preserved by deductive inference. If respecting pretheoretic intuitions is the name of the game, then it would appear that (i) is preferable to (ii) – that accepting (i), in spite of the costs, is the better overall strategy for making our pretheoretic intuitions systematic.

Keith DeRose (1995) defends a version of epistemic contextualism according to which, in a nutshell, it is true to say that one knows that P only if the most similar worlds in which one believes that P falsely via his actual method lie beyond a contextually determined threshold of similarity. According to DeRose, when we assess the knowledge status of a belief, we are obliged to consider possible worlds as dissimilar from actuality as is necessary in order to accommodate the falsity of all contextually salient propositions. It will be true, in a given context, to describe one's belief as knowledge only if one's method gives the right result throughout all of these, and more similar, worlds. DeRose's proposed criteria for contextual salience are quite lenient. In his view, an assertion of the form 'S knows that P' or 'S does not know that P' can suffice to make the proposition P contextually salient.

The framework I have developed here opens up the possibility of a corresponding contextualist account of justification attributions. To develop such a position, one need only substitute 'normalcy' for 'similarity' throughout DeRose's contextualist account of knowledge attributions. It is not my preference to defend

such a view – but it may provide a useful supplement to a DeRose type theory. The primary advantage of DeRose’s theory (as emphasised in his 1995) is its alleged ability to allow for the possibility of knowledge (true knowledge attributions) whilst explaining the appeal of certain sceptical arguments. But there are closely analogous, and equally appealing, sceptical arguments that target the possibility of *justification* rather than the possibility of knowledge (many regard these as constituting the more important and interesting form of sceptical challenge). If we were unable to tell a corresponding story about scepticism of this sort, then surely the theory would be seriously compromised.

My belief that the external world exists not only has an extraordinarily potent default level of safety, it also enjoys an extraordinarily potent default level of reliability. Provided that, in the actual world, there is a reality independent of human minds, worlds in which reality is mind dependent should be regarded not only as an enormous departure from actuality, but also as an enormous departure from ideal normalcy. From the standpoint of the actual world, these worlds should be regarded both as very different and as very strange.

By basing my belief in the existence of the external world upon the appearance of hands before my face, I fail to enhance either the safety or the reliability of the belief beyond the default levels associated with the belief’s content. The appearance of hands before my face will not make a world in which reality is mind dependent any *less* normal than it would otherwise be. That is, the most normal worlds in which reality is mind dependent *and* there appear to be hands before my face will simply be amongst the most normal worlds in which reality is mind

dependent. My account predicts, then, that Moore's inference fails to transmit justification from premises to conclusion. One cannot augment the reliability of a belief in the existence of the external world by holding one's hands before one's face.

IV. MOORE'S 'PROOF'

Tyler Burge (2003) and James Pryor (2004) insist that one could be justified in believing in the existence of the external world in virtue of drawing the inference that Moore recommends. According to Burge and Pryor the only problem with Moore's reasoning is that it would not make for a convincing *argument* – that is, it could never be used to rationally persuade someone who genuinely doubted its conclusion. As far as they are concerned, our aversion to Moore's inference is simply a response to this *dialectical* shortcoming and need not reflect an underlying *epistemic* flaw – such as transmission failure. Whatever the merits of this diagnosis of our aversion, I have proposed a theory of epistemic justification relative to which Moore's inference clearly *is* epistemically flawed. In this final section, I shall offer some reasons for denying that dialectical and epistemic efficacy can come apart in quite the way that Burge and Pryor envisage. I shall focus here upon Pryor's discussion.

Following Pryor, I shall use the term 'doubt' to indicate a deliberate stance toward a proposition – to be distinguished sharply from mere indifference or obliviousness. To doubt that P is true is either to believe that P is false or to regard it as just as likely false as true. It may be that Moore's reasoning could be rationally persuasive for one who has never so much as considered the question of whether there is an external world. After all, it does show me that something I believe – and believe

with a great deal of conviction – commits me to accepting the existence of the external world.

As Pryor is keen to emphasise, any reasoning whatsoever will be rationally unpersuasive for certain doubters. Suppose I have a feeling that Pompey died in the first century AD and, as a result, am dubious of your claim that he died in 48 BC. Suppose you argue as follows, in the hope of persuading me:

- (1) According to Encyclopaedia Britannica, Pompey died in 48 B.C.
- (2) Encyclopaedia Britannica would not get that sort of thing wrong.
- (3) Therefore, Pompey died in 48 B.C.

This argument may or may not prove rationally persuasive. Suppose that I also doubt, for some reason, that encyclopaedias are accurate sources of information. As long as this doubt persists, it would be irrational for me to relinquish my doubts about Pompey dying in 48 BC on the strength of your argument.

According to Pryor, the fact that your argument is unable to rationally move me – with my strange background doubts – does not reflect poorly upon the reasoning. Further, it certainly doesn't prevent *you* from following this reasoning in order to earn a justification for believing the conclusion. This claim is surely correct.

The situation with Moore's argument, however, is somewhat different. Not only would this argument fail to rationally persuade *some* doubters – it would fail to rationally persuade *any* doubter, irrespective of what else he happens to believe or doubt. No one who doubted the existence of the external world could be rationally

moved by Moore's argument. Pryor seems to think that this is not a significant difference. I am inclined to think that it is (as is Jackson, 1987, section 6.4).

If we are out to rationally persuade someone of the truth of a proposition *P*, then sometimes simply asserting that *P* is true will do the trick. It is often rational to accept the testimony of others. Deductive argument, however, can add something to bald assertion. Presenting a deductive argument in favour of *P* involves making a choice. One's audience can glean useful information about the *basis* of one's belief that *P* from the premise choice that one makes (see Jackson, 1987, section 6.2). From your presenting the above argument, I would assume that you have consulted Encyclopaedia Britannica on the matter of Pompey's death and that what you found provides the basis of your belief that he died in 48 BC. Presenting a deductive argument in favour of a proposition *P* is something of a risk, then. It may or may not increase the chances of persuasion – it simply depends upon the nature of one's argument and the background doubts of one's intended audience. Sometimes the less your audience knows, the better.

When one selects a particular deductive argument in favour of a proposition *P*, one effectively advertises a certain basis for believing *P* as 'up for borrowing' by one's audience – to use the phrase favoured by Jackson (1987, chap. 6). When one presents a deductive argument in good faith, one offers one's audience the opportunity to adopt, derivatively, one's own basis for believing *P*. Certain audiences, of course, will spurn the offer.

The framework that I developed in the previous section may be of some use at this point. Suppose we travel as far from ideal normalcy as is necessary in order to accommodate the truth of a putative basis *B* and the falsity of each and every proposition that I doubt (there may not be any single possible world in which *all* of the propositions that I doubt are false). Suppose further that, in none of the worlds subsumed in this sweep is there a world in which *B* is true and *P* false. That is, suppose that in none of these worlds does *B* lead me astray with respect to *P*. This, I suggest, should be regarded a *sufficient* condition for *B* to qualify as a rationally persuasive ground upon which to believe that *P*, for me.

If, of course, one harbours doubts about necessary truths, then there won't be any set of possible worlds that accommodate the falsity of *each and every* proposition that he doubts. For present purposes, though, we can put the possibility of such doubts aside. Often accommodating a subset of a person's doubts – or even just a single doubt – is sufficient to destabilise the modal connection between a proposition and an offered basis upon which it might be believed.

If I doubt that encyclopaedias are reliable sources of information, then the fact that (B) Encyclopaedia Britannica says that Pompey died in 48 B.C., will not qualify as a rationally persuasive ground upon which to believe that (P) Pompey died in 48 B.C., for me. Plausibly, at some of the most normal worlds in which encyclopaedias are not accurate sources of information, *B* is true even though *P* is false. If we travel far enough from ideal normalcy in order to incorporate such worlds, then we effectively disrupt the modal relationship between *B* and *P*. Given, however, that reputable encyclopaedias *are* accurate sources of information in all maximally normal

worlds, the fact that the provision of B could not rationally persuade me of P does not carry any implications as to whether or not it could serve as a contributing reliable basis upon which to believe that P.

Things are different, however, if the provision of a basis B could not rationally persuade *anyone* who doubted that P. When evaluating the rational persuasiveness of B for anyone who doubts P, we are obliged to travel only as far from ideal normalcy as is required in order to accommodate the falsity of P. Doubt of P is, after all, just what all possible P-doubters have in common. If B is a rationally unpersuasive basis upon which to believe that P for any such doubter, then it must be the case that this class of worlds *already* contains worlds at which B is true and P is false. But this, of course, *does* entail that B is not a contributing reliable basis upon which to believe that P. If a deductive argument is, in principle, rationally unpersuasive for anyone who holds doubts about the conclusion, then the basis provided for believing the premises is at best a non-contributing reliable basis for believing the conclusion. In this case, the inference will be an instance of transmission failure.

No external world sceptic could be rationally persuaded by Moore's argument, irrespective of what else they happen to doubt or believe. This serves to show what is, I think, obvious in its own right – namely, once we travel far enough from ideal normalcy to incorporate worlds in which reality is mind dependent, we have already incorporated worlds in which reality is mind dependent and one is having an experience as of hands. It follows from this, in turn, that an experience as of hands could not be a contributing reliable basis upon which to believe in the existence of the external world. The experience of hands before one's face cannot provide one with a

justification for believing in the existence of the external world. Pryor and Burge concede too much to Moorean common-sense epistemology.

Wittgenstein famously described *the external world exists* as a ‘hinge’ proposition, the truth of which is fused into the very foundations of our practice of rational, critical inquiry (Wittgenstein, 1969, §83, 558). Wittgenstein insisted, notoriously, that while such hinge propositions are *certain* for us, we could never have *justification* for believing them. Part of the idea behind this move, I take it, is to synthesise aspects of external world scepticism and aspects of Moorean common sense epistemology, thereby tempering the excesses of both standpoints.

Though this may be an admirable ambition, I do not think that Wittgenstein quite pulls it off. As I have argued, the idea that we could never have justification for believing in the existence of the external world comes at a high price – particularly with regard to the failure of the preservation of justification across deductively valid inferences. The concession that the existence of the external world is ‘certain’ for us is, I fear, insufficient recompense.

If there is a kernel of truth to be found within external world scepticism it is this: We could never have rationally compelling reasons to offer, either to others or to ourselves, in support of the existence of the external world. Neither experience nor ratiocination could provide us with a contributing reliable basis for believing that the external world is there. If there is a kernel of truth to be found within Moorean common sense epistemology it is that we can nevertheless *have* epistemic justification for our conviction. To believe in the existence of the external world is not mere

speculation or whimsy – and the status of the belief is not epistemically fragile. In one sense the very opposite is true – the epistemic strength or security attaching to this belief is consummate.

REFERENCES

- Alston, W. (1988) 'An internalist externalism' *Synthese* v74, pp265-283
- Beebe, H. (2001) 'Transfer of warrant, begging the question and semantic externalism' *Philosophical Quarterly* v51, pp356-374
- Bergmann, M. (2005) 'Epistemic circularity: Malignant and benign' *Philosophy and Phenomenological Research* v69, pp709-727
- Brandom, R. (1983) 'Asserting' *Noûs* v17, pp637-650
- Burge, T. (2003) 'Replies' in Frápolli, M. and Romero, E. eds. *Meaning, Basic Knowledge and Mind: Essays on Tyler Burge* (CSLI Publications)
- Chase, J. (2004) 'Indicator reliabilism' *Philosophy and Phenomenological Research* v69, pp115-137
- Davies, M. (1998) 'Externalism, architecturalism and epistemic warrant' in Wright, C., Smith, M. and Macdonald, C., eds. *Knowing Our Own Minds: Essays in Self-Knowledge* (Oxford: Oxford University Press), pp321-361
- Davies, M. (2000) 'Externalism and armchair knowledge' in Boghossian, P. and Peacocke, C. eds. *New Essays on the A Priori* (Oxford: Oxford University Press) pp 384-414
- Davies, M. (2003a) 'The problem of armchair knowledge' in Nuccetelli, S. ed. *New Essays on Semantic Externalism and Self Knowledge* (Cambridge, MA: MIT Press)

- Davies, M. (2003b) 'Externalism, self-knowledge and transmission of warrant' in Frápolli, M. and Romero, E. eds. *Meaning, Basic Knowledge and Mind: Essays on Tyler Burge* (CSLI Publications)
- Davies, M. (2004) 'Epistemic entitlement, warrant transmission and easy knowledge' *Aristotelian Society Supplement* v78, pp213-245
- DeRose, K. (1995) 'How to solve the sceptical problem' *Philosophical Review* v104, pp1-52
- DeRose, K. (2004) 'Sosa, safety, sensitivity and skeptical hypotheses' in Greco, J. ed. *Ernest Sosa and His Critics* (Oxford: Blackwell)
- Dretske, F. (1970) 'Epistemic operators' *Journal of Philosophy* v67, pp1007-1023
- Dretske, F. (1971) 'Conclusive reasons' *Australasian Journal of Philosophy* v49, pp1-22
- Dretske, F. (2000) 'Entitlement: Epistemic rights without epistemic duties?' *Philosophy and Phenomenological Research* v60, pp591-606
- Dretske, F. (2005) 'The case against closure' in Steup, M. and Sosa, E. eds. *Contemporary Debates in Epistemology* (Oxford: Blackwell Publishing)
- Goldman, A. (1979) 'What is justified belief?' in Pappas, G. ed. *Justification and Knowledge* (Dordrecht: Reidel)
- Goldman, A. (1986) *Epistemology and Cognition* (Cambridge, MA: Harvard University Press)
- Grice, P. (1989) 'Logic and conversation' in *Studies in the Way of Words* (Cambridge, MA: Harvard University Press)
- Heller, M. (1999) 'Relevant alternatives and closure' *Australasian Journal of Philosophy* v77, pp196-208
- Jackson, F. (1987) *Conditionals* (Oxford: Blackwell Publishers)

- Leite, A. (2004) 'On justifying and being justified' in Sosa, E. and Villaneuva, E. eds.
Philosophical Issues 14: Epistemology (Oxford: Blackwell Publishing)
- Lewis, D. (1973) *Counterfactuals* (Oxford: Blackwell)
- Moore, G. (1939) 'Proof of an external world' *Proceedings of the British Academy*
v25, pp273-300
- Nozick, R. (1981) *Philosophical Explanations* (Cambridge, MA: Harvard University Press)
- Pritchard, D. (2002) 'Resurrecting the Moorean response to the sceptic' *International Journal of Philosophical Studies* v10, pp283-307
- Pritchard, D. (2005) 'Scepticism, epistemic luck and epistemic *angst*' *Australasian Journal of Philosophy* v83, pp185-205
- Pryor, J. (2004) 'What is wrong with Moore's argument?' *Philosophical Perspectives*
v18 (Atascadero: Ridgeview Publishing)
- Schmitt, F. (1992) *Knowledge and Belief* (London: Routledge)
- Silins, N. (2005) 'Transmission failure failure' *Philosophical Studies* v126, pp71-102
- Smith, M. (2007) 'Ceteris paribus conditionals and comparative normalcy', *Journal of Philosophical Logic*, v36(1), pp97-121
- Smith, M. (forthcoming) 'What else justification could be' *Noûs*
- Sosa, E. (1999a) 'How to defeat opposition to Moore' in Tomberlin, J. ed.
Philosophical Perspectives v13 (Atascadero: Ridgeview Publishing co.)
- Sosa, E. (1999b) 'How must knowledge be modally related to what is known?'
Philosophical Topics v26, pp373-384
- Sosa, E. (2000) 'Contextualism and scepticism' in Sosa, E. and Villanueva, E. eds.
Philosophical Issues v10 (Boston: Blackwell)

- Sosa, E. (2002) 'Tracking, competence and knowledge' in Moser, P. ed. *Oxford Handbook of Epistemology* (Oxford: Oxford University Press)
- Williamson, T. (2000) *Knowledge and its Limits* (Oxford: Oxford University Press)
- Wittgenstein, L. (1969) *On Certainty*, trans. by Paul, D. and Anscombe, G. (Oxford: Basil Blackwell)
- Wright, C. (1985) 'Facts and certainty' *Proceedings of the British Academy* v71, pp 429-472
- Wright, C. (2000) 'Cogency and question-begging: Some reflections on McKinsey's paradox and Putnam's proof' in Sosa, E. and Villanueva, E. ed. *Philosophical Issues* v10 (Boston: Blackwell)
- Wright, C. (2002) '(Anti-)sceptics, simple and subtle' *Philosophy and Phenomenological Research* v65, pp330-348
- Wright, C. (2003) 'Some reflections on the acquisition of warrant by inference' in Nuccetelli, S. ed. *New Essays on Semantic Externalism and Self Knowledge* (Cambridge, MA: MIT Press)
- Wright, C. (2004) 'Warrant for nothing (and foundations for free)?' *Aristotelian Society Supplement* v78, pp167-212