

現代英米倫理学の批判的検討（上）

—— 超越論的目的論の立場から ——

永井俊哉

本稿は『一橋研究』97号で既に表明した超越論的目的論／目的論的倫理学の立場から、ムーア・ヘアー・サール・ロールズ等の現代英米の、つまり言語分析的倫理学の諸学説を検討する。このことを通して、超越論的目的論／目的論的倫理学の言語哲学的地平を開くことが本稿の目的である。

第一節 ムーア其自然主義的誤謬批判

英米系倫理学に対する大方の常識は次のようなものであろう：論理実証主義的な意味の検証理論は価値述語に真偽の基準を否定し、日常言語学派も規範倫理学を断念し、価値語の意味＝使用の記述に甘んじるメタ倫理学を提唱した。メタ倫理学は道德心理学や道德社会学などの実証科学とは一線を画すけれども、事実と価値を峻別する価値中立的・実証主義的な道德言語学であることには変わりはない。—— こういう常識からすれば、メタ倫理学の開祖たるG.E.ムーアは当然規範倫理学を放棄する端緒を作ったはずだが、その結果はともかくとしても、彼の『倫理学原理』の意図はむしろ逆に規範倫理学の基礎付けにあった⁽¹⁾。彼自身「私は「学問として現れうるであろう全ての将来の倫理学へのプロレゴメナ」を書こうとした」⁽²⁾と序文で明記している。

「倫理学においては他のあらゆる哲学の部門と同じく、その歴史の中にはふんだんに見出される異論や見解の相違は、大抵非常に単純な原因に因るように思われる。その原因とは即ち、答えたいと望む問いがまさにどんな問いであるかをまず知ることをしないで、問いというものに答えようと試みることである」⁽³⁾。ムーアによると「何がよいのか？」という問いで使用されている「よい」という概念には、「目的として善い」と「手段として良い」の二つの異なっ

た意味があって、両者を区別することがまずもって重要である。しばしば（典型的にはカントにおいては）倫理学の本来の問いは「何をすべきか？」であると考えられているが、すべき行為とは目的として善いものを結果として産出する行為であるから、この問いは、さらに基礎的な問いを前提する以上、最も基礎的な問いとは言えない。道徳的行為は「結果を度外視して」なさねばならないとも考えられようが、このような《目的として善い行為》は、《目的として善いもの》という類の一種と考えておく。かくして倫理学における基礎的な問いは「何が目的として善いのか？」であることになるが、「善いもの the good thing」は複合観念であり、「善い good」という単純観念の理解を前提にしている。それゆえ倫理学における最も基礎的な問いは「“善い” とは何か？」「“善い” はいかに定義されるか？」である。この問いが「十分に理解されず、その正しい答えが明瞭に認識されないかぎり、倫理学の残りの部分〔メタ倫理学ではない規範倫理学〕は、体系的知識という観点からは無用も同然である」⁽⁴⁾。

ところがこの問いに対するムーアの解答は「もし“善いとは何か？”と問われるならば、私の答えは“善いとは善いである”であって、それで終わりである。またもしも“善いはいかに定義されるべきか？”と問われるならば、私の答えは、“それは定義されえない”であって、私がそれについて言わなければならないことはそれが全てである」⁽⁵⁾という一見「極めて失望させるもの」であった。ではなぜ「善い」は定義できないのだろうか？「善い」が定義できないならば、彼が目指す「体系的学問（science）としての倫理学」⁽⁶⁾など不可能なはずである。ムーアは「善い」を定義することは自然主義的誤謬であるというが、なぜいかなる点でそれは《誤謬》であるのか？これに対する彼の説明は必ずしも明解ではなく、複数の解釈が可能である。そこで以下それらを類型化しつつ順次検討してみよう。

1. まず「善い」は「善い」であって、他とは意味が異なり、この異なるものを混同することが自然主義的誤謬であるという「定義」を他の言葉で置換しうる」の意に採った解釈が考えられる。ムーアはあたかもこの解釈を仄めかすかの如く『倫理学原理』の冒頭に「全てはそれがあるところのものであって、それ以外のものではない」というバトラーの標語を掲げている。しかしこの解釈によると、「善い」が例えば「快い」で定義できないのは「善い」が「快い」

と違う言葉だからである、ということになるが、「この意味で“善い”が定義できないというわけではない」⁽⁷⁾。事実ムーアは、「善い」が「内在的価値」「存在すべき」と同じ意味であると言って、前者を後者によって“定義”しているのである⁽⁸⁾。ムーアによると、定義には（１）私が「馬」と言う時、それは「馬属の有蹄四足獣」を意味するものとする、というような恣意的な（規約主義的な？）言語的定義、（２）人々が「馬」と言う時、それは「馬属の有蹄四足獣」を意味している、というような本来の（クワイン的な？）言語的定義、つまり経験社会学的・辞書の定義、以上のような言語的定義とは別に、（３）「馬」とは四足・頭・心臓・肝臓等々が一定の関係で並んでいるものであるというように複合物を部分的に分解して、その配列を記述する（ラッセル流の？）実在的定義があるが、彼が問題にしているのは（３）の意味での定義なのである。

2. そこで自然主義的誤謬とは「善い」は単純で分析不可能であり、これを複合物として定義しようとするものである、と解釈される⁽⁹⁾。「善い」は、単純で部分を持たないがゆえに定義ができない⁽¹⁰⁾。「ちょうど“黄色い”が単純観念であるように“善い”は単純観念である。つまりちょうど諸君が、それをあらかじめ知らない人〔例えば色盲者〕に黄色が何であるかをいかなる方法によっても説明することができないように、諸君は善いは何であるかを説明することはできないのである」⁽¹¹⁾。ところがムーアが挙げている自然主義的誤謬の例には、「善い」を「欲せんと欲すること that which we desire to desire」のような複合物で定義している場合のみならず、「善い」を「快い」のような単純観念で定義している場合もある⁽¹²⁾。それゆえ少なくとも自然主義的誤謬の説明としてはこの解釈は成り立たない。ムーア自身『倫理学原理』第二版序言の草稿（1920—1年頃・未完）の中で、「善い」が分析不可能であるから定義できないという議論が誤っていたことを認めている⁽¹³⁾。むしろ重要なのは「善い」が「独特の対象 a unique object」⁽¹⁴⁾であり、複合的か単純であるかを問わず、「善い」とは違った種類の観念と同一視されてはならない、ということであろう。ムーアが「善いについての命題は全て総合的であって、決して分析的ではない」⁽¹⁵⁾と言うのはこの意味においてである。本来総合判断である、つまり主語と述語が異質である「Xは善い」が「善いとはXであり、かつその時のみである」という分析判断に転化されるとき、自然主義的誤謬が犯される。

そこで問題は、「善い」がいかなる点で「独特」で、他の観念とは「違った種類」であるのか、ということになるのだが、ここですぐ思い付きそうなのは、

3. 「善い」は非自然的属性であって自然的属性によっては定義されえず、自然的属性によってこれを定義することが自然主義的誤謬であるという解釈である⁽⁶⁾。このように定義すれば、確かになぜ「善い」を定義する誤謬が“自然主義的”誤謬であるかのかが明確に成る。「もし[快いなどと]同じ意味で自然的対象ではない“善い”を、およそ何らかの自然的対象と混乱するならば、それを自然主義的誤謬と呼ぶ理由がある」⁽⁷⁾。ところが他方ムーアは、「善い」を「完全性」「永遠性」「超感性的」などの非自然的属性と同一視する形而上学的倫理学をも“自然主義的”誤謬を犯したとして、自然主義的倫理学の時と全く同じ論証形式で論駁している。それゆえ「自然主義的誤謬」は必ずしも適切な名称とは言えないのだが、「我々がその誤謬に出合った時それを認識するならば、それをどう呼ぶかはどうでもよい問題なのである」⁽⁸⁾。いずれにせよ「善い」の特異性はその非自然的性格にあるのではない。そこで次に、

4. 「善い」は倫理的属性であって非倫理的属性によっては定義されえず、これを定義することが自然主義的誤謬であるという解釈が考えられる⁽⁹⁾。このように解釈すれば、なぜ「善い」が「正しい」や「望ましい desirable」などで定義されうるが、「快い」や「欲求されている desired」では定義できないかが明確と成るが、同時にこの二分法はヒューム以来の《Is》と《Ought》、つまり理性と道徳感情の区別を彷彿させる⁽¹⁰⁾。実際ムーアが「“実在はこのような性格である”と主張する何らかの命題から、“これはそれ自身において善い”と主張する何らかの命題を我々が推論できる、または確証することができると考えることは、自然主義的誤謬を犯すことである」⁽¹¹⁾と言う時、あたかも自然主義的誤謬とは事実から価値を導出することであるかのようだが、蓋しこれが最も一般に流布している解釈なのである⁽¹²⁾。

この解釈の最大の難点は、なぜ自然主義的誤謬が誤謬であるのかがあまり判然としないところにある。なるほど非倫理的な命題から倫理的な命題を導出することは、両者を混同して、後者を前者に還元してしまう（あるいは前者を後者に還元してしまう）ことになるのだから誤謬であるという説明がなされるかもしれない⁽¹³⁾。しかしそれは言語的誤謬であっても、倫理的誤謬ではない。ムーアの関心はたんにメタ倫理学にあったのではなく、規範倫理を基礎付ける

かぎりでのメタ倫理学にあったことを想起しなければならない。例えば、自然淘汰により生物は進化して来たという事実命題から、ゆえに弱肉強食の生存競争はあるべきだという規範命題を導出するソーシャル・ダーウィニズムを我々が批判するとき、果たして我々は当為が存在に還元されていること、結論における「べし」が“情動的意味”を欠いていることを批判しているのだろうか？

断じて否。情動的意味が欠けているなら、むしろ誰もソーシャル・ダーウィニズムを批判しないぐらいである。我々が批判するのは、推論内容が誤っているから、即ち自然淘汰が常に善いとはかぎらないからである。逆に云えば「善くない（もちろんなぜ善くないかが争点と成るのだが）自然淘汰」という“反証例”（例えば植民地での圧政、弱者の大量殺戮等々）を提示しないかぎり、そしてなぜそれが悪いかを正当化しないかぎり、我々はソーシャル・ダーウィニズムを批判できないのである。このことを倫理学以外の例で見てみよう。「白鳥は白い」という命題は、もし白鳥が白さそのものであるならば、それは「白鳥であるとき、かつそのときのみ白い」という同語反復になる。あるいはそこまで言わなくても、「白鳥」が“swan”ではなくて“white bird”であるならば、当の命題は分析的に真となる。そしてこのような誤謬がおかされると、「白鳥が白色か否か」という動物学的探求が無意味になる。しかしここで注意しなければならないのは、「白鳥」と「白色」が異質であり、前者から後者が分析的に導出されないからといって「白鳥は白い」という判断が直ちに不可能になるのではなくて、むしろ「白鳥が白色か否か」という動物学的探求を有意味にするということである。「白鳥は白い」という総合判断を論駁するためには、白色でない白鳥（例えばオーストラリアに棲息するクロハクチョウ）を反証例として提示しなければならない。より高度な自然科学の命題ではさらに反証例の解釈が問題となるであろうが、それは倫理学でも事情は同じである。以上から明らかなように、事実と価値は存在性格が異なることは、両者を主語と述語として結合した判断が偽であるための必要条件であっても十分条件ではなく、したがって倫理学が有意味であるための不可欠条件なのである。倫理学は、事実から価値を正しく導くことができるかもしれないのであって、ムーアはそれが《直覚》によって（積極的にではなくあくまでも消極的にではあるが）できると考えていた。そこで結局4.の解釈は、

5. 自然主義的誤謬とは、「善いもの」という複合観念における単純観念

「善い」以外の非倫理的属性Pを「善い」と混同して同一視し、Pが唯一の善であると考えてしまう誤謬である⁽⁹⁰⁾という解釈に移行する。すなわち、Pが極めて頻繁に「善い」と同時に生起するとき、Pと「善い」との間に結合が生じ、Pが「善い」の代わりに用いられるようになる。この虚偽を暴露するためには、Pであっても善くはない反証例を挙げればよい。例えば「自然は善い」という自然主義的倫理学の命題においては、自然という善いものが持つ「正常な」という属性と「善い」という属性との間に結合が生じて、しばしば「正常な」 $\cap$ 「善い」の結合は、「正常な」 $\equiv$ 「善い」にまで進展する。しかし後者はもちろん前者も成り立たないことは、 $\sim$ 「正常な」 $\cdot$ 「善い」（例えば、ソクラテスやシェークスピアなどの“異常な”天才）を反証例として提示することによって暴露される⁽⁹⁵⁾。形而上学的倫理学を批判するときも形式は同じである。例えば「最高善は永遠の实在である」と言う時、「超感性的」と「善い」が混同されている。「超感性的」という点では「悪のイデア」も考えられるはずであり、これが反証例となるのだが、この種のイデアの存在を認めようとしなかったプラトンは、広い意味での「自然主義的」誤謬を犯していたわけである。

ここでムーアが批判していることはPを唯一の善と見做すことであって、Pがある一つの善いものであると判断することまでが誤謬であるとは言っていないことに注意しなければならない。定義できないのは「善い」であって「善いもの」は《定義》できるのである：「もし私が“善い”だけでなく“善いもの”までが定義できないと考えているならば、私は倫理学の本を書かなかっただろう。というのも私の主たる目的は、善いものの定義を発見することを助けることなのだから。いま私が善いを定義できないということを強調しているのは、そうした方が“善いもの”の定義を求めるときに誤る危険が少なくなるであろうと考えるからに他ならない」⁽⁹⁶⁾。ムーアは、この「善いもの」の定義は《直覚》に拠るしかないと考えて、直覚主義的倫理学を打ち立てようとする。——このムーアの議論は意外である。というのも、彼の自然主義的誤謬批判は、自然主義のみならず直覚主義をも批判するだけの潜在能力を持っているからである。つまり $(\exists x)(Px \cdot \sim Gx)$ という反証例は、 $(\forall x)(Px \equiv Gx)$ という自然主義者の分析的命題のみならず、 $(\forall x)(Px \supset Gx)$ という直覚主義者の総合的命題をも否定するからである⁽⁹⁷⁾。反証例があるにもかかわらず

ず、そのような $P \sim G$ は「真の」 P ではないとうそぶきつつ、「 P が唯一の善なり」と主張することが自然主義的誤謬であるとするならば、「 P はいつでもどこでも誰にとっても善きものなり」と主張する誤りは直覚主義的誤謬とも名付けられよう。ムーアは、その実体主義的絶対主義的発想⁽²⁸⁾ のゆえにこの誤謬を犯した。だが我々は、彼の批判の射程を拡大することによって、当の彼の規範倫理学をも破一壊することができると考える。しかしムーア批判を云々する前に、彼の《目的として善いもの》即ち「内在的価値 intrinsic value」についての議論を一瞥しておく必要がある。

ムーアは、一方では内在的価値は直覚によってしか知りえないという無媒介性のテーゼを掲げながらも、他方ではその内容確定に際しては彼なりに方法を工夫していた。「絶対的孤立の方法 method of absolute solution」⁽²⁹⁾ がそれである。「[何が内在的価値を持つかという] この問いに関して正しい解決に到達するためには、何が絶対的孤立においてそれ自身で存在していても、我々がその存在を善いと判断するであろうようなものであるかを考察する必要がある」⁽³⁰⁾。

この方法は快樂主義批判に際して駆使されるのだが、実は以前の直覚主義的誤謬批判の時と同様に、 $\sim (\forall x) (P x \supset G x)$ と論理的に等値な $(\exists x) (P x \cdot \sim G x)$ が示されるだけなのである。ムーアは『ピレボス』(21A)を引用しつつ次のように主張する：もしも快が唯一の善であるならば、快以外の何物も、例えば記憶も意識も知性もなくとも善いはずなのだが、そうだとするならば、自分がかつて快を感じていたことも・今快を感じていることも・将来快を感じるであろうこともわからなくなり、かくして牡蠣のごとき生活を送らなければならないだろうというわけである⁽³¹⁾。快と快の意識が同じではないことは確かだとしても、色の無い青色がないように、意識の無い快はありえないのではないのか、とムーアの読者は反論したくなるであろう⁽³²⁾。しかし一歩譲って、快樂主義者が唯一の善と見做すものが快の意識であると認めたとしても、具体的意識内容を欠いた抽象的な快の意識を絶対的に孤立化した場合、それが唯一善いものであるかどうかは極めて疑わしい⁽³³⁾。我々が「善い」と判断するとき、そこに常に快の意識が随伴しているからといって、快の意識が「善いもの」であるとか、況んや「善い」そのものであるなどとは言えないのである。

全ての事象を意識作用に還元する認識論的主観主義は、新実在論者ムーアが常に否定する立場であった。対象と作用の区別は『観念論論駁』に詳しい⁽³⁴⁾。ムーアによると、あらゆる観念論の根底には「実在するとは知覚されることである」 $(\forall x)(Essex \supset Percipix)$ というパークリー以来の命題が横たわっているが、これを否定するのに「知覚されていない実在物」 $\sim (\exists x)(Essex \cdot \sim Percipix)$ という反証例を「考える」ことはできない。なぜなら「考える」ことも広義の「知覚する Percipere」(つまり認識作用)である以上、その実在物は依然として Percipiであるからである。しかし両者は実在的に区別されないにしても概念的には区別されうであろう。例えば「青の知覚」は「青に対する知覚」であって、知覚自体は青くない。たとえ「青の知覚」と「青い知覚」が同じであっても、我々はなお「認識の対象」と「対象の認識」を区別することができるのである⁽³⁵⁾。同様に青色が好きな人の場合であるが、「青の快」は「青に対する快」であって、快自体は青くない。たとえ「青の快」と「青い快」が同じであっても、我々はなお「快の対象」と「対象の快」を区別することができるのである。快楽主義者は快の対象から快の意識を孤立化させることができなかったで、「善い」と「快の意識」を同一視したのである。

ムーアはこの経験的意識に当てはまることを超越論的意識についてまで当てはめようとする。ムーアによると、カントの理論哲学においては「正しい」を「一定の方法で認識されている」と同一視し、実践哲学においては「善い」を「一定の方法で意志されている」と同一視したが、この《コペルニクス的転回》も対象と作用の(ノエマとノエシスの)混同なのである。「Ich denke は、全ての私の表象に伴いなければならない」をもじって言えば、「Ich will は、全ての私の格率に伴いなければならない」という前提に立って、カントはあの形式的倫理学を打ち立てたのである。かくしてカントは、「善い」を「善き意志」＝「定言命法」、すなわち「ある実在的で超感性的な権威によって命令されている」で定義する自然主義的誤謬を犯したというわけである⁽³⁶⁾。

以上から明らかなように、「絶対的孤立の方法」は以前の反証例提示の方法と目的と形式は同じである。但し、反証例提示の方法では、混同されている非倫理的属性Pが特定の対象であったので、反証例を一つ提示すればそれで足りたのだが、全ての善の判断に随伴する作用がPである場合は、概念上これを孤立化させて、その内在的価値を吟味する次善の方法が採られたわけである。こ

の絶対的孤立の方法は、反証例提示の方法とは異なって、全ての自然主義的／直覚主義的誤謬を暴露する包括的な方法なので、内在的価値をこの方法によって検証された価値として定義することができる。ではいかなる内在的価値が絶対的孤立においてもなお多くの価値を保持しうるのだろうか？ ムーアは明言していないがおそらく皆無であろう。彼自身は「人間間の交際の楽しみ」と「美的対象の享受」を内在的善と考えている⁽³⁷⁾が、例えば後者などもそれを（１）対象に備わる美的性質（２）これに感動しうる情緒性（３）対象が実在することの正しい信念という各構成要素に分解し、それらを単独で考察した場合、ほとんど価値がないことになってしまう⁽³⁸⁾。最も価値がありそうな（１）の美的対象、例えば美しい絵画も、額・縁・画布・絵の具・油等々の部分に分解すれば、どれも価値がないことが判明する。（２）や（３）に至っては、もしその対象が醜いものならば、積極的に悪くさえある。

この事態に気が付いたムーアは、自分の立場を守るべく「有機的統一体の原理 principle of organic unities」なるものを提唱する。すなわち「全体の内在的価値は、その部分の価値の総計と同じでもなければ比例もしない」⁽³⁹⁾。論敵の快樂主義を攻撃するときには全体を部分へと分解してその価値を貶めておきながら、自分の立場を説明する段階になると、部分には還元されえない全体の価値を云々する彼のやり方は一見いかにも卑怯に見える。しかしここで彼が否定しているのは自然主義的誤謬であって直覚主義的誤謬ではないということ想起しなければならない。彼は、快が内在的価値の一構成要素であることを認めるに吝かではないのであって、ただそれが唯一の内在的価値を僭称することを批判しているだけなのである。「この〔有機的統一体の原理を無視するという〕誤謬は、もし全体の一部分が内在的価値を持たないならば、全体の価値は全て他の部分に存しなければならないと考えられるとき犯される。かくして、もし全ての価値ある全体が、たった一つの共通属性しか持っていないと見做されうるならば、この属性を所有しているからというだけで全体は価値があるに違いないと通常考えられてきた。そして当の共通属性は、それだけで考えられると、それだけで考えられたそのような全体の他の諸部分より大きい価値を持っていると思われるなら、その幻想は大いに強められる。しかし、我々が当の属性を孤立して考察して、その属性をそれを部分とする全体と比較すれば、当の属性が持っている価値は、それだけで存在しているときその属する全体が持つ

ている価値には遠く及ばない、ということがたやすく明らかになるだろう」⁽⁴⁰⁾。ここから以前の自然主義的誤謬の5の解釈を次のように表現することができる：自然主義的誤謬とは有機的全体としての内在的価値を特定の非倫理的な部分の価値に還元する誤謬である。

ムーアによれば、内在的価値の全体を構成する部分は、それ自体無価値であるどころか積極的に悪くさえあることもある⁽⁴¹⁾。だがこのように価値を、全体が持つ創発的特性 (emergent property) として特徴付けることによって、彼は不可知論に陥ってしまったのではないか？ 部分を事実、全体を価値に置き替えたとき、部分に対する全体の独立性のテーゼは、実は事実から価値が導けないことの婉曲な表現であることがわかるであろう。ムーアの規範倫理学が不可知論に終わった原因は、しばしばそう考えられているように、彼のメタ倫理学的な問題設定にあるのではなくて（この反省的＝超越論的問題設定自体は我々も積極的に受け継がなければならないであろう）、彼の価値概念の実体主義的・絶対主義的性格にあるのではないのか？

ムーアによると「ある種の価値が“内在的”であるということは、あるものがそれを所有するかどうか、そしてどの程度それがそれを所有するかという問いが、当のものの内在的本性にのみ依存するということを意味することに他ならない」⁽⁴²⁾。内在的価値は、超時間空間的不変性と超個体の普遍性を持って内在的本性に依存する。この内在的価値の概念規定から、彼が実体的な内在的価値の絶対的妥当性を信じていたと考えることができるであろう。ムーアによれば、関係概念としての価値は“外在的”な手段としての価値に過ぎない。すなわち、内在的価値とは「性質 (character) であって、関係的属性 (relational property) ではない」⁽⁴³⁾。しかしその反面ムーアは「善い」「美しい」などの述語が第一性質ではないことはもちろん「黄色い」のような第二性質とも存在性格を異にした述語であることをも了解していた：「黄色さと美しさは、両方ともそれを所有しているものの内在的本性にのみ依存する述語であるが、黄色さがそれ自身内在的述語であるのに対して、美しさの方はそうではない」⁽⁴⁴⁾。では価値のこの特殊な存在性格とは何か？ ムーアによれば、価値とは部分に還元されない全体であった。この全体とは部分と部分の関係であると言えないか？ 次のように考えてみよう：

〔テーゼ1〕 価値とは事象に独立自存する実体はなく、主体（但し経験的な

主体）と客体との関係である：この実体概念と関係概念の対比は、カッシーラーのそれである。価値とは、主体と客体を包括する欲求なる（疑似）志向性における、心的契機と物的契機の関数値である。世界一内一存在する主体が客体に対して関係を持つ（sich verhalten）ということは、主体が客体に対して態度を採る（sich verhalten）ということである。主体との関係を絶たれた《絶一対 absolute》の孤立における客体は、理論的客観的に観察される「対一象 Gegen-stand＝主体に対して立つもの」として価値を失う。もちろんここで言っている《関係》は、レアールな《人と物との関係》であって、イデアールな《人と人との関係》ではない。経験的な主体と客体との関係は、超越論的反省において概念的に述定されなければならない。そこで次に「よい（善い＋良い）」の意味（基準ではない！）を次のように定義しよう：

〔テーゼ2〕「よい」とは、我々実践主体が持つ目的に対する手段／形態の有効性である：我々は、ムーアの「目的として内在的に善い」と「手段として外在的に良い」の区別を撤廃し、「～にとって／としてよい」で統一することができる。①「～にとってよい」と②「～としてよい」という我々の区別は、どちらも①目標的目的—実現手段、②理念的的目的—実現形態の関係を表す概念であって、ムーアの区別とは異なる。このテーゼ2から、直覚主義的誤謬を、ある価値が目的との関係を離れて無制約的に妥当すると考えられた時に犯される誤謬と定義することができる。この誤謬が犯されたとき、PがGと結合して、あたかもGはPという「内在的本性にのみ依存する」かの如く思念されてしまうのである。

この我々の目的論のテーゼに対して、はたして目的は手段を正当化しうるのか、と疑問を持つ向きもあろう。しかしそれは、目的の概念を狭く採ることによって生じる⁽⁴⁵⁾。例えばある過激学生が、革命の目的のためには手段を選ばないと豪語したとする。ところが彼は、革命の目的と手段の転倒であることに気が付いて、もっと穏健な手段を選ぼうとするようになるかもしれない。その場合確かに革命という目的が全ての手段を正当化するわけではないが、正当化しないことを正当化しているのはより高次の目的であることに注意しなければならない。要するに《悪しき目的》の反価値性は、それを手段／形態とするより高次の目的に対する不適合性なのである。目的—手段／形態の系列を上昇して行けば、究極目的に到達するであろうが、この究極目的（最高諸目的の体系

の実現の形式＝超越論的主観性）自体は善でも悪でもない。ちょうど物体の総体に重さがないように価値の総体には価値がないのである。

究極目的は空虚な統一性であって内容を持たない。内容を持った理念的目的のうちで最も抽象的な＝最高の目的は、究極目的によって統一される諸目的である。最高諸目的は、究極目的の（１）一つの（２）実現形態である。（１）は自然主義的誤謬に陥らないための、（２）は直覚主義的誤謬に陥らないための要件である。すなわち究極目的の実現形態のありかたは、ある個人においてそうあるよりも他でありうるし、他の諸目的との関係で、ある最高目的の実現に（per seにではないがper accidensには）価値がない場合もありうるのである。では究極目的の内容はどのようにして決めるのだろうか？ この問いに対しては次のように答えたい：[テーゼ3] 目的から手段／形態を選ぶ前に、手段／形態から目的の内容を確定しなければならない。ムーアは具体的な手段／形態の価値を捨象していきなり「目的として善いもの」を捜そうとしたために直覚主義、ひいては不可知論に陥った⁽⁴⁶⁾。我々はこれに対して現存する全ての価値を出発点にしつつ、その根拠を一なるものへと収斂させ、そこから現存の行為規範・制度を正当化／修正する。そしてそれが我々が以前提唱した目的論的還元・構成・破壊であったわけである。

第二節 ヘアーの普遍的指令主義

R.M.ヘアーは自分の立場を「普遍的指令主義universal prescriptivism」⁽⁴⁷⁾と呼んでいる。つまり彼のメタ倫理学上の主張は、情動主義批判としての普遍主義と記述主義（自然主義＋直覚主義）批判としての指令主義批判から成り立っているわけなのだが、ムーアとの関連上、普遍化可能性を取り上げる前に、まず後者の側面から見ていこう。“Principia Ethica”以後、「善い」の定義可能性をめぐる様々な議論がなされたが、エイヤーは、「善い」の定義が不可能なのは「善い」のような倫理的概念が疑似概念（pseudou-concept, 真偽の判定の不可能な概念）だからであると考えた⁽⁴⁸⁾。つまり倫理的概念には、道徳的感情を表出したり、聞き手の行動を刺激したりする情動的な機能があり、それゆえそれは事実に還元することができないというわけである。これは自然主義的誤謬の4の解釈に相当する⁽⁴⁹⁾。エイヤーの説を受け継いだスティーブンソンも「これはよい」の意味を「私はこれを是認する、君も是認したまえ」

すなわち、感情表出+聴者説得というように分析している⁽⁵⁰⁾。

ところでこれらは「善い」の《意味》の定義であって《基準》の定義ではない。情動主義者は「善い」とそれが適用される対象との関係をどのように考えていたのだろうか？ 彼等は、前者を情動的意味、後者を記述的意味と名付けて両意味契機の相互独立性を強調する。ピューラーが分析したように、言語には表出・喚起・叙述の三機能がある⁽⁵¹⁾ののだが、情動的(e-motive)意味の“情”(感情表出: express)と“動”(行動喚起: move)の機能、記述的意味の事象叙述の機能がそれぞれ当の三つに対応する。例えば「虎だ!」という発話は、虎出現を叙述すると同時に話者の恐怖と当惑の感情を表現し、聴者に逃走の行動を喚起するが、同じ発話を加藤清正がするとき、それは獲物発見の喜びの表現であり、部下への準備の指令である、というように記述的意味は特定の情動的意味と必然的な結合関係を持たない。同様に特定の記述的意味は、他人の特定の情動的意味と必然的に結合しているわけではない。それゆえある特定の支持理由(supporting reason)を挙げることによって、常に他人を自分が持つ特定の態度に同調させることができるとは限らない⁽⁵²⁾。合理的方法によってもなお態度の不一致が残存する場合、通常我々は、非合理的な説得的方法(persuasive method)、つまり様々な物理的心理的方法で相手の感情に因果的な影響を及ぼし、説得させようとする⁽⁵³⁾のであって、これがスティーブンソンの心理学的分析の材料となる。この種のメタ倫理学においては、規範倫理学への道が全く塞がれていることはいうまでもない。

さてヘアーは、正当化と動機付け、発語内的行為(illocutionary act)と発語媒介的行為(perlocutionary act)との相違、即ち「誰かにあることをするように言う過程が、彼にそれをさせることとは相互に論理的に全く異なっている」⁽⁵⁴⁾ことを指摘して、因果分析から概念分析へ、情動主義から指令主義へとメタ倫理学の方向を転じる。彼は、情動的意味の感情表出よりも行動喚起の機能を重視し、指令的意味と記述的意味の相互関係の分析をおのれの課題とする。

しかしながらヘアーは、一方ではこのように情動主義を批判しておきながら、他方「善い」の意味に関しては情動主義的な定義をしている。彼が「自然主義の諸理論における誤謬は、それらが価値判断を事実言明から導出することができるようにしようとするあまり、その判断にある指令的推薦の要素を無視することにある」⁽⁵⁵⁾と言うとき、明らかに自然主義的誤謬を4の意味で解釈してし

まっている。このような解釈をする人は、《ヒュームの法則》を信奉して、事実と価値／意味と基準を峻別し、結局は価値主観主義＝不可知論的相対主義に陥らざるをえない仕組みになっているのである。ヘアーは「よい」の意味を「選ばれるべき」や「勧めるべき」で定義するが、このような意味を知ったところで「何がよいのか」の問題は解決されえないどころか、解決の糸口すら与えられないままである。かくして「よい」を対象に適用するときの基準の判定は「そのつどの新しいレッスン」⁽⁵⁶⁾になってしまう。対象にどのような価値を付着させるかは、個人の恣意的な決断、即ち《自由》に任せられているのである。

とはいえ、規範倫理学にも取り組む以上、価値問題に関して全く無関心でいるわけにはいかないのであって、それゆえ初期のヘアーは、「よい」を「べし」で定義することによって、つまり価値問題を実践問題に還元することによって事態を解決しようとした。「AはよいXである」(Aは類Xの一つの種)は「Aは他のXの種Bよりもよい(better) Xである」を意味し、これは「もしXを選ぶべきなら、他の種よりAを選ぶべきだ」を意味するという定義がそれである⁽⁵⁷⁾。ところが中期のヘアーは、このような定義が結局《趣味の押し付け》につながることに気が付き、「AはBより善い人だ」から指令「BはもっとAのようになろうとすべきだ」は導出されえないと言って、初期のテーゼを撤回している⁽⁵⁸⁾。彼のリベラリズムは、後期においては選好(preference)の内容に対する不干渉という形で保持されているのだが、実はこの価値自由主義が彼の功利主義をかなり平板なものにしてしまうのである。しかしこの側面からヘアーを攻撃する前に、彼の普遍的指令主義のもう一つの側面である普遍化可能性(universalizability)の議論を見ておく必要がある。

ヘアーはまず様相論理学風に命令文のwhatとhow、指示句(phrastic)と断定辞(neustic)を区別してその合理的側面を確保し、これを基にして実践論理学を樹立しようとする。即ち「君はドアを閉めようとしている」という平叙文と「ドアをしめたまえ」という命令文は、「君が近い未来にドアを閉めること」という指示句を共有しており、両者の相違は断定辞の違いに過ぎず、まさにこの指示句を共有するがゆえに実践論理学にも通常の論理学と同じ論理規則(例えば矛盾律)が成り立つのである⁽⁵⁹⁾。だから「ドアを閉めたまえ、しかしドアを閉めるな」という命令を下すことはできない。この一見トリヴィアルに見

える論理的事実が、実は道徳推論の基礎となるのである。

さて普遍的命令文は、二人称の人物の未来のある行為を命じているのだから普遍的ではないのだが、記述的な指示句を持つがゆえに、少なくとも論理的には常に普遍化可能である。即ち「道徳法則は、記述的判断が普遍化できるのと全く同じ仕方で普遍化できるが、これは道徳表現にも記述的表現にも記述的意味があると言う事実から来ている」⁽⁶⁰⁾。換言するならば、全ての個別的な命令は、それをone exampleとして包摂する普遍的原則を暗に前提しているということなのである。

では普遍化という概念操作は、具体的にはどのようなものなのだろうか？ヘアーは、「厳密に言うならば…普遍化に異なった諸段階があるというわけではないということを強調しておきたい。道徳的判断はただ次の意味においてのみ普遍化可能である、すなわち道徳的判断は、その普遍的属性という点で同じである全ての場合に対して同じ判断を導出する、これが私の主張である」⁽⁶¹⁾と言うが、実際には普遍化は、全称化と普遍化の二つのプロセスから成り立っている⁽⁶²⁾。両者の違いを知るためには、そもそも普遍的でない命題にはどのような種類があるのかを確認しておく必要がある。まず差し当って次のような三つの命題を考えてみよう：

- ①「嘘をついてはいけない場合も時にはある。」
- ②「永井俊哉は常に嘘をついてはならない。」
- ③「日本人は常に嘘をついてはいけない。」

この道徳判断の記述的意味（指示句）の部分だけ記号化すると次のようになる（但し、Wは「言葉である」Mは「ひとである」D「～が～を公言する」Lは「嘘である」Jは「日本人である」を表す記述記号、nは「永井俊哉」を表す個体定項とする）：

- ① $(\exists x)(\exists y)(Wx \cdot (My \supset Dyx) \cdot \sim Lx)$
- ② $(\forall x)(Wx \cdot Dnx \supset \sim Lx)$
- ③ $(\forall x)(\forall y)(Jx \cdot Wy \cdot Dxy \supset \sim Lx)$

ここから上掲の各判断がなぜ普遍的原則ではないのかは明らかである。①は存在量化が冠頭にあるので特称判断であり、普遍的でない以前にそもそも原則

ではない。②と③は全称判断であるから原則にはなっているのだが、②は n という個体定項を含んでいるので個別的な全称判断であり、③は個体定項を含んでいないので、そのかぎりでは非個別的であるが、「日本人である」という普遍的ではない述語を含んでいるので、結局両方とも普遍的な原則ではないのである。ところで②と③は異なるものであろうか？ そうではあるまい。N「永井俊哉である」という述語記号を作れば、②は $(\forall x)(\forall y)(N_y \supset D_{yx}) \supset \sim L_x$ と書き替えられるし、N「 $\sim$ は $\sim$ の国籍をもつ」という述語記号を作り、「日本」を表す個体定項 j を作ってやれば、 J_x は N_{xj} と書き替えることができる、というように②と③は相互に変換可能だからである。そこでこの両者を合わせて個別的な全称判断とし、特称判断からこれへの移行を全称化、これから普遍的な判断への移行を普遍化と呼ぶことにしよう。

全称化という概念操作は、実は我々が第一節で直覚主義的誤謬を暴露するために用いた方法と同じものである。つまり $(\exists x)(P_x \cdot \sim G_x)$ という価値経験に出合わないかぎり、私念 $(\exists x)(P_x \cdot G_x)$ は $(\forall x)(P_x \supset G_x)$ へと全称化可能である。しかしこの概念操作によっては利己主義は論駁されえない。例えば P が「永井の利益になる」を表すなら、エゴイスト「永井」はこの全称価値判断を喜んで無条件的に受け入れることができるであろう。そこで道徳的推論においては全称化可能性のみならず、普遍化可能性までもが要求されることになる。普遍化可能性に関して中期のヘアーが挙げている例は、 A が B から、 B は C から借金をして、 B が A を借金を取り立てるために投獄しようか否かと迷っている場合である⁽⁴⁰⁾。いま B が利己的な動機から A を投獄しようと決心したとする。この時“Let me put A into prison !”という個別的指令を、その理由を条件法導入して記号化すると、次のようになる： $\sim Fab \supset Gba$ (F 「 $\sim$ は $\sim$ から借金を返済している」 G 「 $\sim$ は $\sim$ を投獄する」)。 B はこの個別的指令を普遍化した $(\forall x)(\forall y)(\sim F_{xy} \supset G_{yx})$ という原則を受け入れることになるが、彼はまさにこのことによって、自分が $\sim F_{bc} \supset G_{cb}$ というもう一つの個別的指令にも言質を与えている (commit oneself) ことに気が付くのである。かくして B は、 G_{yx} を欲すると同時に欲しないという《意志の矛盾》に陥るのだが、もし彼の preference が $Gba < \sim Gcb$ であるならば、 A の投獄を断念すべきだ、という結論を出すことになるであろう。

このような普遍化の構成要素として、ヘアーは1. $\sim Fab$, $\sim F_{bc}$ という事

実、2. 個別的指令を普遍化する論理、3. $G_{ba} < G_{cb}$ という嗜好、の三つを差し当り列举する⁽⁶⁰⁾が、1や2だけでなく3が必要であることに注意しなければならない。当事者の嗜好を知ることは一種の事実認識なのであるが、判断に指令性を与えるものとして格別の位置が与えられているわけである。ヘアーは、この三つにさらに4. 想像力という要素を加える。つまり、たとえCに相当する人物が居なくても、なお想像力によってそれを仮想して、ないし(同じことだが)想像力によって相手(A)の立場に自分を移し置いて、“果たして自分は G_{ba} の指令を意志することができるだろうか?”とBは自問しなければならない。この4を追加することによって、普遍化可能性のテストは利己的打算的なものから道徳的なものへと変貌する。道徳的推論において重要な役割を果たすこの《他者の立場へ「私」を移し置く》という思惟の運動が可能であるのは、「私」が「非本質的属性non-essential property」だからであると後期のヘアーは言う⁽⁶¹⁾。それゆえ「もしもジョーンズがスミスならば…」という仮定は、両者がそれぞれ別のessenceを持っているのだからナンセンスであるが、「もしも私がスミスであるならば…」という仮定は有意味である、と彼は説明するのだが、要するに「ジョーンズ」や「スミス」などの固体定項を、人格一般に適用可能な「私」へと普遍化することによって、道徳的共感は論理的に可能となるということなのである。

以上の1～4の要件を満たしながらも、なおBは、自分とAとのある相違を指摘することによって、 $G_{ba} \sim G_{cb}$ という指令を可能にするかもしれない。例えばAは独身で、貸与された金を自分の娯楽のために浪費していたが、Bには扶養すべき家族がいて、然るべき事情のため返金できないでいるということにしておこう。この種の情状酌量の対象となる「然るべき事情」があることをHとすると、Bが依拠する普遍的原則は、 $(\forall x) (\forall y) (\sim Hx \cdot \sim Fxy \supset Gyx)$ というように修正される。このように次々と例外が条件法導入によって繰り入れられて行くなれば、その原則はもはや普遍的でなくなるのではないかと懸念する人は、普遍と一般を混同しているのである。「『一般的 general』の反対は『限定的 specific』であり、『普遍的 universal』の反対は『個別的 singular』である」⁽⁶²⁾。 $(\forall x) (\forall y) (\sim Fxy \supset Gyx)$ という一般的原则も $(\forall x) (\forall y) (\sim Hx \cdot \sim Fxy \supset Gyx)$ という限定的原則も、共に普遍的原则なのであって、 $(\exists x) (\exists y) (\sim Fxy \supset Gyx)$ や $(\forall a) (\forall b) (\sim$

$H a \cdot \sim F a b \supset G b a$) などの個別的原則からは区別される。一般的か限定的かは程度の差なのであるが、普遍的か個別的かはそうではない⁽⁶⁷⁾。一般的普遍が限定的となるプロセスを自然科学の対応例で見てみよう。科学者は、27°C 5気圧のもとでの1モルの窒素の気体の体積は5リットルであるという個別的な気体の記述を気体の状態方程式 $P V = n R T$ へと普遍化する。ところがこの“一般”法則は、“理想”気体の状態方程式と言われるように、実在的気体の状態に適合しない場合（低温・高圧で凝縮する場合）がある。そこで科学者は「もし気体分子に大きさがなく、分子間力が働かないとすれば」という前提を条件法導入することによってこの法則を“限定的”にするが、このことは法則をあいまいにするのではなくて、むしろ厳密にするのであって科学的認識の進歩の一部である。同様に道徳原則の場合も「諸種の例外を認める際、我々がしていることは、原則をあいまいにすることではなくて、厳密にする」⁽⁶⁸⁾ ことなのであって、それは「我々の道徳的発展の一部なのである」⁽⁶⁹⁾。

この条件法導入に際して、導入される条件が正当であるか否かをどうやって決定するのかという問題が生じてくる。また借金の例に戻るが、例えばBは、Aが黒人でBが白人であるという事実を指摘して、これがmorally relevantであると主張しだすかもしれない。「ある状況のある特徴を道徳的に関係があると見なすことは、その特徴に言及している道徳原則をその状況に適用することである」⁽⁷⁰⁾。そこで適用されている道徳原則は $(\forall x)(\forall y)(\sim (W x \cdot B y) \cdot \sim F x y \supset G y x)$ ということになる ($W = \text{White}$, $B = \text{Black}$)。ところがBがこの原則を信奉するのは、それが白人である自分の利益につながるからなのであって、このような利己主義者に対するヘアーの説得方法は、相変わず想像力を働かせることである。もしもBが、自分と相手の皮膚の色が逆転した想像上の場合でも、なおこの原則を信奉し続けることができるだろうか？と自問すれば、彼は皮膚の色がmorally irrelevant propertyであることを認めざるをえなくなるであろう。

我々は、以前普遍化において示されるのは、行為原則の自己矛盾ではなくて、普遍化された原則と嗜好（つまり選好）との衝突であることに留意しておいた。しかしここで衝突しているのは選好というよりもむしろ選好を満足させることを指令する原則であるから、要するにこれは原則相互の、借金の例で言うならば、復讐欲の原則と護身欲の原則との葛藤であると考えられる。このように直

覚レヴェルでさしあたり妥当する二つの可能的原則（prima facie principles）がある状況で葛藤するとき、批判レヴェルでどちらの原則を優先させる（override）かを選好の強度という点で比較考量して決定することが功利主義的倫理学の仕事である、とヘアーは考える。この比較考量に際して、ヘアーは1. その時の2. その人の選好を十分に想像する（fully represent to oneself）ことを強調する⁽⁷¹⁾。換言するならば、今の自分の選好でその時のないし他人の選好に関することから決定してはいけないのである。例えば老人医療費は税金の無駄であるという発言は、将来の自分のことを考えると撤回を余儀なくされるであろう。このような批判的思考をし続けることによって、我々は結果的には《最大多数の最大幸福》を実現することができるかもしれない。だがヘアーの療法的な功利主義は、従来の建設的な功利主義と同じではない。「我々が我々の方法に要求するのは選好の強度の比較である。我々は快や幸福の単位を総計する必要はない」⁽⁷²⁾。ヘアーは現存する直覺的な道德的原則を全面的に肯定しているのであって、ただこれらが葛藤を起こしたときのみ調停に乗り出すのである。彼の功利主義がイギリスに伝統的なリベラリズムの立場に立つものであることは、もはや明らかであろう。

ヘアーのこのリベラリズムの遠因は、以前仄めかしておいたように、その情動主義的な自然主義的誤謬解釈、そしてそこから生じてくる価値主観主義・不可知論的相対主義、ここにある。「善い」の意味と基準のつながりを見出すことができなかった彼は実践哲学において、多様な選好内容を度外視して、たんに外面的な形式的側面しか扱えないはめになっている。そこでまず彼の「よい」の定義から疑ってかかればならない（続）。

註

- (1) ムーア自身はmeta-ethics なる語を使っていない。岩崎武雄氏によれば W. K. Frankenaの発案のようである。『岩崎武雄著作集第三巻』280頁参照。そのフランケナもまた、倫理的タームの意味のメタ倫理学的分析の目的は規範倫理学のjustification であると位置付けている。

cf. his Ethics, Chapter 6. p. 95. Vgl. auch H. Albert ; Ethik und Meta-Ethik, Das Dilemma der analytischen Moralphilosophie, in : Werturteilsstreit, Darmstadt 1971, Hrsg. E. Topitsch, SS. 173-175. ただしアルバートはメタ倫理学の根本動機は規範倫理学に対

する中立性のテーゼであると考えている。

- (2) G. E. Moore ; *Principia Ethica*, Cambridge University Press, 1903, p. ix.
- (3) *Principia Ethica*, p. vii.
- (4) *Principia Ethica*, p. 5.
- (5) *Principia Ethica*, p. 6.
- (6) *ibid.*
- (7) *Principia Ethica*, p. 8.
- (8) *Principia Ethica*, p. 7.
- (9) A. C. Ewing, *Ethics*, Free Press 1953, p. 79.
- (10) *Principia Ethica*, p. 9.
- (11) *Principia Ethica*, p. 7.
- (12) cf. *Principia Ethica*, §.13.
- (13) C. Lewy ; G. E. Moore on the Naturalistic Fallacy, in : G. E. Moore *Essays in Retrospect*, ed. A. Ambrose and M. Lazerowitz 1970, pp. 293f.
- (14) *Principia Ethica*, p. 16.
- (15) *Principia Ethica*, p. 7.
- (16) C. D. Broad ; Certain Features in Moores Ethical Doctrines, in : *The Philosophy of G. E. Moore*, ed. P. A. Schilpp Cambridge 1942, pp. 57-67.
- (17) *Principia Ethica*, p. 13.
- (18) *Principia Ethica*, p. 14.
- (19) Frankena, *The Naturalistic Fallacy*, in : *Readings in Ethical Theory*, ed. W. Sellars and J. Hospers 1970, p. 56. and also *Ethics*, p. 99.
- (20) 周知の《ヒュームの法則》をヒューム自身が本当に主張していたかどうかは疑問である。彼は「[isからought への変化は] 極めて重要である。というもこの“べし” ないし “べきでない” はある新しい関係または断定を表現しているので、それを注視し・説明し、同時に、この新しい関係がどのようにして全く異なったものから導出されるのかという極めて不可解に思われることに理由が与えられることが必要だからである」と言っており、isからought を導出することができないとは言っていない。
D. Hume ; *A Treatise of Human Nature*, p. 469.
- (21) *Principia Ethica*, p. 114.
- (22) 日本のものでは、岩崎武雄『現代英米の倫理学』勁草書房1963年(62頁)などがそうである。
- (23) アルバートは、自然主義的誤謬推論を1. 道徳的命題の記述的誤解(SollenのSeinへの還元)と2. 記述命題から規範命題を導出する誤謬推理

(SeinのSollenへの還元) という相互に正当化し合う二つの謬見を合わせたものとして説明している (Albert ; op. cit. SS. 482f.). 自然主義的誤謬推論の問題は自然主義的誤謬判断の問題である :

大前提	進化 は 善い
小前提	自然淘汰 は 進化である
結 論	故に 自然淘汰は 善い

事実判断 (小前提) から価値判断 (結論) が導出できるかどうかは結局大前提の判断が妥当か、つまり事実を表す主語に価値の述語を付けられるかどうかに係っている。ムーアが問題とする自然主義的誤謬とは、総合判断である倫理的判断が分析判断化されることであるが、その分析判断には次の二つが考えられる。1. 「事実=事実」: この場合「進化は善い」は「進化は進化である」と同じ意味になる。つまりは「進化≡善い」という定義がなされているのだが、この「善い」にはなんら「指令の意味」がない。これが自然主義的誤謬だというのが、ヘーアの「自然主義的 (記述主義的) 誤謬」の解釈である。2. 「価値=価値」: この分析判断化の場合、「進化は善い」の主語「進化」(あるいはより適切には「進歩」) の中にあらかじめ価値が潜んでいて、この判断はそれをただ取り出しただけになる。この命題を主張している人には「進化」のもたらす悪しき側面を反証例として提示しても「それは進化ではない、衰退だ」と答えるので無駄である。

——実証主義者は「事実からは価値は出て来ない」と言い張り、解釈学者 (と仮にこう名付けておくが) はこれに対して「裸の事実などありえない。事実は常に価値と共にあるのだから、価値判断は可能である」と応酬する。実証主義者曰く「事実は事実だ!」解釈学者答えて曰く「価値は価値だ!」——ムーアはこのような不毛な対立地平そのものを克服しようとしたのである。

(24) cf. Principia Ethica, p. 10.

(25) Principia Ethica, § 27.

(26) Principia Ethica, pp. 8f.

(27) この証明 (註: P=nonethical property, G=good as end)

1. (1) $(\exists x) (Px \cdot \sim Gx)$ 仮定

2. (2) $PY \cdot \sim GY$ 仮定

3. (3) $(\forall x) (Px \supset Gx)$ 仮定

2. (4) PY (2) $\cdot -$

2. (5) $\sim GY$ (2) $\cdot -$

3. (6) $PY \supset GY$ (3) $\forall -$

2, 3. (7) GY (4) (6) $\supset -$

2, 3. (8) $GY \cdot \sim GY$ (5) (7) $\cdot +$

2. (9) $\sim (\forall x) (Px \supset Gx)$ (3) (8) $\sim +$
 1. (10) $\sim (\forall x) (Px \supset Gx)$ (1) (2) (9) $\exists -$

証明は省略するがこの逆の導出も成り立つので、

$$(\exists x) (Px \cdot \sim Gx) \equiv \sim (\forall x) (Px \supset Gx).$$

(28) 我々は「実体主義」と「偶有主義(?)」の対立地平に対して「関係主義」の立場を、「絶対主義」と「相対主義」の対立地平に対して「相関主義」の立場を採る。ムーアはラッセルと共に、ブラッドリーを始めとする当時のイギリスのヘーゲリアンに対抗して new realism の旗手となっていた哲学者で、ラッセルと同様その哲学は(そして当然その倫理学も)素朴実在論的原子論的性格が強い。

(29) Principia Ethica, p. 188.

(30) Principia Ethica, p. 187.

(31) Principia Ethica, pp. 88f.

(32) A Duncan-Jones, Intrinsic Value : Some Comments on the work of G. E. Moore, in Essays in retrospect, pp. 315-318.

(33) Principia Ethica, p. 91.

(34) G. E. Moore ; The Refutation of Idealism, in : Philosophical Studies, Routledge & Kegan Paul 1922. この論文は "Principia Ethica" と同じく 1903 年に発表された。

(35) ibid, p. 26.

(36) このムーアのカント批判が不当であることについては, H. J. Paton The Categorical Imperative, p. 43.

(37) Principia Ethica, p. 188.

(38) Principia Ethica, p. 199.

(39) Principia Ethica, p. 184.

(40) Principia Ethica, p. 187f.

(41) Principia Ethica, §§. 129-133.

(42) Moore ; The Conception of Intrinsic Value, in : The Philosophical Studies, p. 260.

(43) Moore ; Is Goodness a Quality?, in : Philosophical Papers, George Allen & Unwin 1959, p. 97.

(44) Moore ; The Conception of Intrinsic Value, p. 272.

(45) ムーア自身「正しい」[手段として良い]は「善い結果の原因」以外の何物をも意味しないし、意味しえず、したがって「有用な」と同一である」(Principia Ethica, p. 147) と言って, the end will never justify the means という常識を否定している。

(46) F. Kaulbach, Ethik und Metaethik, Darstellung und Kritik metaethischer Argumente, Wissenschaftliche Buchgesellschaft

- 1974, SS. 79f/82f.
- (47) Hare ; *Moral Thinking, Its Levels, Method, and Point*, Oxford 1981, p. 228.
- (48) A. J. Ayer ; *Language, Truth and Logic*, Penguin Books 1936, p. 142.
- (49) エイヤーはその実証主義的解釈から当然ムーアの直覚主義を批判することになる。ムーアの「神秘的な知的直観」(ibid, p. 140) が客観的妥当性を持たないので、エイヤーはそれを主観的感情にまで矮小化したわけである。
- (50) C. L. Stevenson, *Ethics and Language*, Yale University Press 1944, p. 21. スティーブンソンは「君も是認したまえ」の部分は心理的状态の記述で心理学の対象となり、そのためある程度の合理性を持つと言っている (ibid, p. 26) が、これは当たっていない。エイヤーが言うように、表出は「何ら命題を表現しているのではない」(Ayer, op. cit, p. 142)。それは陳述の内容 (Was) を語っているのではなくて、様式 (Wie) を示している。
- (51) K. Bühler : *Sprachtheorie, die Darstellungsfunktion der Sprache*, 1934, S. 24.
- (52) Stevenson, *Ethics and Language*, p. 30.
- (53) Stevenson, *Ethics and Language*, p. 139.
- (54) Hare ; *Language of Morals*, Oxford 1952, p. 13.
- (55) *Language of Morals*, p. 82.
- (56) *Language of Morals*, p. 102.
- (57) *Language of Morals*, p. 184.
- (58) Hare ; *Freedom and Reason*, Oxford 1963, p. 155.
- (59) *Language of Morals*, p. 17f.
- (60) *Language of Morals*, p. 30.
- (61) *Moral Thinking*, p. 108.
- (62) 全称化可能性については、内井惣七氏の「倫理的判断の普遍化可能性について」(人文学報38号1974年)を参照。内井氏のこの論文には次の三つの点で疑問を感じた：1. 氏は「冠頭標準形が全称記号で始まるけれども個体定項を含んでいる」命題は個別的の全称形を持ち、「冠頭標準形が全称記号で始まり、かつ個体定項を全く含まない」命題は非個別的の全称形を持つと言う(同22頁)。しかし本論で示したように、両者の間には本質的な相違はない。飯田隆氏を援用して言えば、個体定項「ソクラテス」は「ソクラテスるもの」というように述語+変項に置換できる(同氏『言語哲学大全I』劉草書房1987年、208頁)。内井氏は「非個別的の全称判断を前提している」を全称化可能性と呼んでいる(24頁)が、「非個別的」は削除されるべきではないかと思う。2. 氏はまた「記述的意味を持つ判断が全

て普遍化可能性であるというのは誤りである」(25頁)とヘアーを批判する。例えば「ニクソンはアメリカ人である」という記述的判断は、「ニクソンと適切な点で類似の人は全てアメリカ人である」と全称化するが、この適切な点とはアメリカ国籍を持つということであり、これはアメリカという個体に言及する性質であるので普遍化されえないというわけである。しかしこの全称判断は「その国の国籍を持つ人はその国の人である」というように普遍化されるのではないか？ 非道徳的行為は、その普遍化を当人が欲求できないという意味で普遍化不可能なのであって、論理的には全ての原則は普遍化可能なのである。内井氏は、倫理的判断は記述的意味を持つがゆえに普遍化可能なのではないというテーゼから、さらに普遍化可能性の規範的性格を主張する(29頁)が、普遍化可能性そのものはたんに論理的可能性として事実的に存立しており、現実には普遍化して自分の原則を吟味することが各人の規範的当為なのではないだろうか？ 3. 「法的な当為はその範囲内では全称化可能であるが普遍化可能ではない。例えば、同じ売春という行為でも、日本では現在これを法的に禁じているけれども、西ドイツの一部あるいはアメリカのネヴァダ州の一部では合法的である(空間的な制約)。また同じ日本においても、売春禁止法が発効する以前と以後とは話が違ってくる(時間的な制約)」(32頁)。この法と倫理の区別にも問題がある。《実定法は、成立時から廃止されるまでかつその立法府の管轄領域内においてのみ妥当である》ことを条件法導入すれば、法もヘアーが言う意味においてuniversalでありうるからである。

- (63) Freedom and Reason, 6. 3.
- (64) Freedom and Reason, p. 93.
- (65) Moral Thinking, p. 119f.
- (66) Freedom and Reason, p. 39.
- (67) Freedom and Reason, p. 40.
- (68) Language of Morals, p. 52.
- (69) Language of Morals, p. 54.
- (70) Moral Thinking, p. 63.
- (71) Moral Thinking, p. 105.
- (72) Moral Thinking, p. 124.