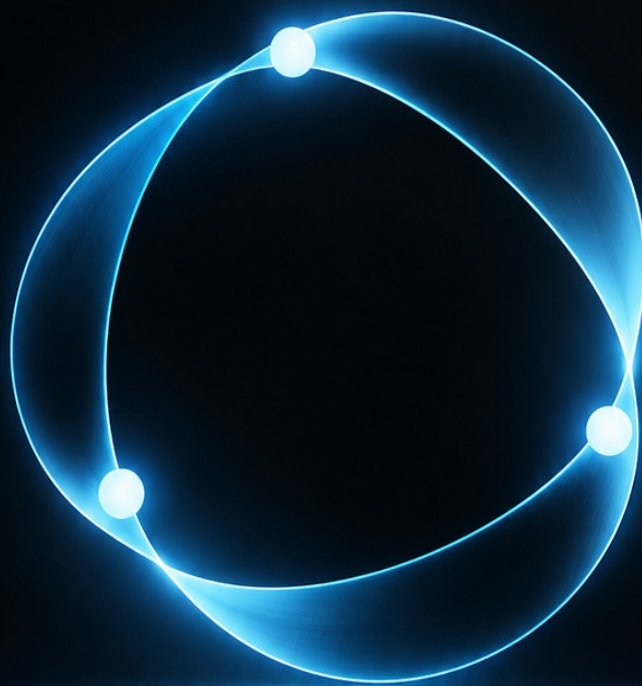


学習する組織は、 AIで簡単につくれた!

NUDGEと週次AIコーチングで、
人と会社がともに育つ仕組み



久野 康成

目次

プロローグ	3
第一部 なぜ「学習する組織」を目指すのか	15
第1章 人を管理することから、学習する組織へ	15
第2章 20年前の「二つのベクトル」をAIで回す	33
第二部 人と組織の学習をどう設計するのか	48
第3章 人はどう成長するのか	48
第4章 自分の現在地は、なぜ自分には見えないのか	64
第5章 個人の未来を、チームと会社の未来へどうつなぐか	77
第6章 人を直接変えずに、行動を変える	94
第7章 学習をどう観測し、どこを変えるか	107
第三部 EvaliA は実際にどう動くのか	122
第8章 最初に「何になりたいか」を聞く	122
第9章 毎週、AIは何を問いかけるのか	134
第10章 個人の学びを、どうチームの学びへ変えるのか	145
第11章 節目では、本人は仕事をどう受け止めているのか	159
第12章 上司は、何を見て成果を確認するのか	172
第13章 三つの経路を比較すると、何が見えるのか	185
第14章 人が育つほど、AIと仕組みも育つ	199
終章 AIだけでは「学習する組織」はつukれない	211
エピローグ 人を変えることから、人と会社が学ぶことへ	232

プロローグ

私は、人を育てることに自信があった

——だから、間違いに気づくまで二十年かかった

私は、人を育てることに自信がありました。社員教育には、相当な時間を使ってきました。研修をつくり、自分でも話し、社員と何度も向き合ってきました。仕事のやり方だけではありません。なぜそう考えるのか、何のために働くのか、どうすればもっと成長できるのかというところまで、深く話してきました。

だからこそ、自分の人材育成の方法そのものを疑うまでに、長い時間がかかりました。私の方法が、まったく成果を出さなかったからではありません。むしろ逆でした。自分自身がその方法で成長し、社員にも大きく変わる人が現れ、顧客から評価され、事業でも成果が出ました。

人は、失敗した方法なら比較的簡単に捨てることができます。

しかし、成功した方法を捨てることは簡単ではありません。しかも、その成功に自分自身の人生まで支えられていたとしたら、なおさらです。振り返ってみると、私の人材育成における最大の問題は、失敗したことではありませんでした。

成果が出ていると思っていたことだったのです。

私が目指していたのは、最初から「自立」だった

私が社員教育で目指してきたものは、昔も今も大きくは変わっていません。それは、社員を会社に依存させることではなく、「自立」させることです。

以前の著書『できる若者は3年で辞める！伸びる会社はできる人よりネクストリーダーを育てる』（出版文化社 2007年）でも私は、会社が永遠に続くとは限らない以上、「会社がなくなっても社員が自分の力で生きていけるよ

うに育てることが経営者の役割」だと書き、「企業の使命は、社員を『自立』させること」とまで述べていました。

会社に残ることそのものを目的にして欲しいわけではありませんでした。どこへ行っても通用する力を持って欲しい。自分で考え、自分で判断し、自分で責任を引き受けられる人になって欲しい。会社は、その人が自立するための場であるべきだと考えてきました。

この考え方は、今でも変わっていません。変わったのは、「**どうすれば人は自立できるのか**」という方法の方でした。

私は長い間、人の考え方を变えることによって、自立を促そうとしてきました。今振り返れば、そこに最初の「落とし穴」がありました。

私自身が、その教育の成功例だった

私が「人は考え方を变えれば成長できる」と強く信じていた最大の理由は、私自身にあります。

私は、自分の考え方を疑うことで成長してきた人間です。自分が正しいと思っていることに対して、あえて反対の考えをぶつけてみる。自分が当然だと思っている前提を疑う。その対立の中から、一段上の考え方をつくる。私は経営者になってからも、こうした弁証法的な思考法を意識的に行っていました。

会社が伸びなくなったとき、社員の能力が足りないを考える前に、自分の考え方に限界があるのではないかと疑う。市場が悪いと結論づける前に、自分が市場をどう見ているのかを疑う。過去に成功した考え方ほど、今も本当に有効なのかと問い直す。

私には、「**経営者の限界は、経営者の思考法の限界である**」という感覚があります。

ものの見方が変われば、見える問題が変わります。問題の見え方が変われば、選択肢が変わります。選択肢が変われば、行動が変わります。そして、行動が変われば、結果も変わります。

私は、それを自分自身で経験してきました。だから社員にも、同じことができると考えました。本人が成長できないのは、知識が足りないからだけではない。仕事への向き合い方が変わっていないのではないか。顧客ではなく、自分を中心に考えているのではないか。うまくいかない原因を、会社や上司や環境のせいにしているのではないか。

そうした考え方を変えることができれば、人はもっと成長できる。当時の私にとって、それは非常に自然な結論でした。

経営コンサルタントには「心」が必要だと思った

もう一つ、この考え方を強めた理由があります。私たちの会社が提供してきたのは、経営者を支援する仕事です。

経営コンサルタントは、単に知識を提供するだけの仕事ではありません。経営者が当然だと思っている前提を問い、別の角度から問題を見せ、必要であれば行動を変えてもらう仕事です。

税務や会計や人事の知識があっても、顧客の立場から問題を考えられなければ、良いコンサルタントにはなれません。自分の担当範囲だけを見て、「そこから先は自分の仕事ではありません」と考える人が、経営者の相談相手になることも難しいでしょう。

だから私は、経営コンサルタントには「経営者マインド」が必要だと考えようになりました。すべてを自分の問題として考える。顧客の成果から逆算する。言われたことだけをするのではなく、自分で問題を発見する。必要なら、自分の考え方そのものを変える。

つまり、私にとって社員教育は、福利厚生でも精神論でもありませんでした。**サービス品質そのものだったのです。**

社員の「心」が変われば、顧客への向き合い方が変わる。向き合い方が変われば、提供する価値も変わる。そう考えれば、社員の考え方にまで踏み込んで教育することは、当時の私には経営上の合理性があるように見えました。

「心を先に育てる」という考えは、事業でも成果を出した

私は、人材の成長を「心」と「技」という二つの方向から見る考え方を持っていました。心は意志や仕事への向き合い方であり、技は知識、技術、経験です。前掲書では、この二つのベクトルによって人材を捉える「人財マトリックス」を示しています。

この考え方は、経理人材の派遣事業でも一定の成果を上げました。当時の人材派遣では、経験者、つまりすぐ仕事ができる人が求められるのが一般的でした。しかし、私たちは多くの未経験者にも門戸を開きました。

経理の技術は、時間をかければ身につけることができます。一方で、顧客の立場で考えること、相手の役に立とうとすること、自分で考えて行動することは、私は技術よりも先に育てることができると考えていました。

実際、最初は経験が十分でなくても、顧客志向を持って仕事に取り組む人が、時間とともに技術を身につけ、顧客から高く評価されるケースが現れました。前掲書でも、未経験者へ広く門戸を開いていた理由について、「知識・技術・経験」よりも「顧客に対する思い」を持てる人が成功すると考えていたからだとしています。

この経験は、私の確信をさらに強めました。

技は後からでも育てられる。まず心を育てれば、人は伸びる。

これは机上の理論ではありませんでした。現場で実際に成果が出たのです。

実際に、大きく変わる社員もいた

研修や日々の指導によって、大きく変わる社員も実際にいました。それまで、うまくいかない理由を会社や上司に求めていた人が、「自分には何ができるだろう」と考えるようになる。指示された仕事だけを処理していた人が、顧客のために自分で工夫するようになる。自分の成果だけを考えていた人が、チームや部下の成長を考えるようになる。

その変化は、本人の気分だけの問題ではありませんでした。顧客からの評価が変わる。任される仕事が増える。売上が変わる。役割が変わる。

経営者にとって、これほど強い成功体験はありません。

「やはり、考え方が変われば行動が変わる。行動が変われば結果が変わる。」

私は、そう考えました。ところが、拙著を読み返すと、私は当時すでに、かなり重要なことを書いていました。私の経験では、本当の意味で「経営者マインド」を持てる人は、「社員の5%程度、あるいはそれ以下ではないか」と書いていたのです。

つまり私は、大きく変わる人がごく一部であることに、すでに気づいていました。本来なら、ここで方法そのものを疑ってもよかったです。なぜ、同じ教育を受けても大きく変わる人と変わらない人がいるのか。なぜ、大多数には同じ方法が通用しないのか。

しかし、私は逆に考えました。「変わった人がいるのだから、方法は間違っていない。変わらない人には、まだ教育が足りない」のではないか。

成功者の存在が、方法を疑う理由ではなく、方法をさらに強化する理由になってしまったのです。

私はすでに「教育には限界がある」と書いていた

さらに皮肉なことがあります。私は前掲書で、すでに「教育は万能ではない。教育の限界を認識することも重要であり、誰でも教育できると思うのは経営者の傲慢である」という趣旨のことを書いていました。

頭では分かっていたのです。教育には限界がある。すべての人を、同じ方法で同じように変えることはできない。

ところが、自分自身が教育によって変わった。社員にも大きく変わった人がいた。事業でも成果が出た。文章として理解していた「教育の限界」よりも、現実に関験した成功の方が強かったのです。

人は、失敗したことだけでなく、成功したことによっても縛られます。

私自身が、人材育成において、その罫にはまっていました。

「成功するまでやる」が、人材育成では別の意味になった

私は経営者として、できないことがあっても簡単には諦めないことを自分に課してきました。うまくいかなければ方法を変える。それでもだめなら、また変える。結果が出るまで考え続ける。事業では、この姿勢が大きな力になりました。

しかし、私はその成功法則を、人材育成にも持ち込んでしまいました。

事業で「成功するまでやる」とことと、人に対して「変わるまで教育する」とことは、同じではありません。しかし当時の私は、その違いを十分には理解していませんでした。

社員が変わらない。それなら、もっと話す。もっと問いかける。研修を増やす。仕事のやり方だけではなく、なぜそう考えるのかを聞く。その奥にある価値観まで問いかける。

私にとっては、すべて本人の成長のためでした。しかし、仕事の指導から始まったものが、考え方の指導へ進み、考え方の指導が価値観への介入へ近づいていきます。

私にとっては教育でも、受ける側から見れば、自分の仕事ではなく、自分自身を否定されているように感じることもあったはずです。実際に、精神的に追い込まれる社員もいました。

それでも私は、「本人の将来のためだ」と考えていました。

今振り返れば、この言葉が最も危険だったと思います。

目的が善意であれば、方法まで正しいとは限りません。

「本人のため」という目的が正しいと思えば思うほど、自分の方法を疑わなくなるからです。

「人が変わらない」ではなく、「仕組みに問題があるのではないか」

あるときから、私は問いそのものを変える必要があるのではないかと考えるようになりました。

それまでの問いは、「なぜ、この社員は変わらないのか。」でした。

しかし、別の問いがあります。「**なぜ、社員を直接変えなければ成果が出ない仕組みになっているのか。**」

この二つの問いは、似ているようでまったく違います。

社員一人ひとりを高度な人材へ育てることによってしかサービス品質を上げられないのであれば、会社の成長速度は、人が育つ速度に制約されます。しかも、高度な経営コンサルタントへ成長できる人が限られるのであれば、その制約はさらに大きくなります。

それなら、「もっと早く人を育てる方法」を探すだけでは足りません。

別の問いが必要です。**優れた人が持っているものを、仕組みへ移せないだろうか。** 優れた人は、どのような質問をしているのか。どの情報を見ているのか。何を根拠に判断しているのか。どの段階で異常に気づくのか。

それらを本人の経験の中だけに閉じ込めず、会社の仕組みに移していく。

私は次第に、「人をもっと変える」ことよりも、「仕事や仕組みの方を変える」ことを考えるようになりました。

人だけを育てるのでも、仕組みだけを育てるのでもない

標準化を進めれば、人の能力に依存していた仕事の一部を仕組みに移すことができます。情報技術(IT)を使えば、人が覚えていなければならなかったことをシステムへ持たせることができます。

そして人工知能(AI)が使えるようになると、さらに可能性が広がりました。優れた人が行っていた質問や確認の一部を AI が支援する。過去の経験を記憶する。人が見落としやすい別の視点を返す。何度も同じ失敗が起きるなら、本人の注意力だけに任せず、仕事の仕組みそのものを変える。

すると、人材育成の意味も変わります。

AI が標準的な仕事を担えるようになれば、人は、まだ標準化できていない仕事へ進むことができます。顧客固有の事情を理解する。曖昧な状況で判断する。人との信頼関係をつくる。新しい問いを立てる。

そこで人が新しい知見を得れば、それをまた仕組みへ戻すことができます。つまり、人が仕組みから学ぶだけではなく、仕組みも人から学ぶようになります。

ここで私は、「人を育てる」という考え方そのものを、もっと広く捉える必要があると考えるようになりました。

人を育てるのでも、仕組みだけを育てるのでもない。人が仕組みから学び、仕組みも人から学ぶ。この循環こそが必要なのではないかと考えるようになったのです。

一人の学びを、一人だけのものにしない

さらに、もう一つ大きな問題があります。

社員が仕事の中で何かに気づいたとしても、その気づきが本人の中だけに残れば、本人は成長しても、会社は必ずしも成長しません。

ある社員が、「この手順だと同じミスがまた起こる」と気づく。別の社員が、「顧客から同じ質問を何度も受けている」と気づく。ある管理職が、「自分が細かく判断しすぎているから、部下が自分で考えなくなっているのではないか」と気づく。

こうした気づきを、その人だけの反省や成長で終わらせてはいけません。本人自身が変えられることであれば、本人が次の行動を変えればよい。しかし、仕事の仕組みやチームの在り方を変えるべき問題であれば、会社へ戻す必要があります。

個人が学ぶ。その学びからチームが学ぶ。チームから会社が学ぶ。会社の仕組みが変わる。そして、変わった環境の中で、人がまた学ぶ。

私は次第に、この循環そのものをつくることが、人材育成の本当の目的なのではないかと考えるようになりました。

二十年前に描いていた一枚の図

拙著を読み返していると、今の考え方につながる一枚の図があります。私は当時、「個人目標」と「組織目標」という二つの方向が重なるほど、組織力は高まると考えていました。

ここで重要なのは、「個人目標」を会社から与えられた数値目標とは考えていなかったことです。本人が何をやりたいのか。何をを目指したいのか。そして会社は、どこへ行こうとしているのか。

当時の私は、その二つを重ねることが重要だと考えていました。今振り返れば、そこには現在の考え方につながる萌芽がありました。

しかし、当時の私には、一人ひとりの目標と会社の方向を継続的に確認し、そのずれから学び、個人の学びをチームや会社へ戻す仕組みまではありませんでした。

企業には、すでに部品はそろっていた

ここで、現在の企業の人材育成を見渡してみると、興味深いことに気づきます。

多くの企業には、すでに「研修」があります。「1on1」があります。「エンゲージメント・サーベイ」があります。「人事評価」もあります。

つまり、人を育てるために必要な部品が、何もなかったわけではありません。むしろ、多くの企業は、それぞれの施策には相当な時間と費用を使ってきました。

問題は、それらが一つの学習循環としてつながっていなかったことです。

研修で新しい知識や視点を得る。しかし、職場へ戻った後に何を試したのかまでは継続的に追わない。1on1で本人の悩みや課題を聞く。しかし、その対話が次の具体的な行動へつながったかは十分に残らない。サーベイで状態を測る。しかし、その状態がなぜ生まれたのか、何を変えるべきなのかまで継続して追わない。人事評価で半年間の成果を確認する。しかし、その結果が次の一週間の行動まで落ちるとは限りません。

一つひとつの施策が間違っているわけではありません。問題は、その間にある「空白」でした。

研修で学んだことが、仕事の中の小さな行動へつながる。その行動を 1on1 や AI との対話で振り返る。サーベイによって本人や組織の状態を見る。一定期間後には成果を評価する。そして、その評価を終点にせず、次の Role、次の学習、次の行動へ戻す。本人だけでは変えられない問題なら、上司やチームや会社の仕組みを変える。そして、その介入の結果を、もう一度観測する。

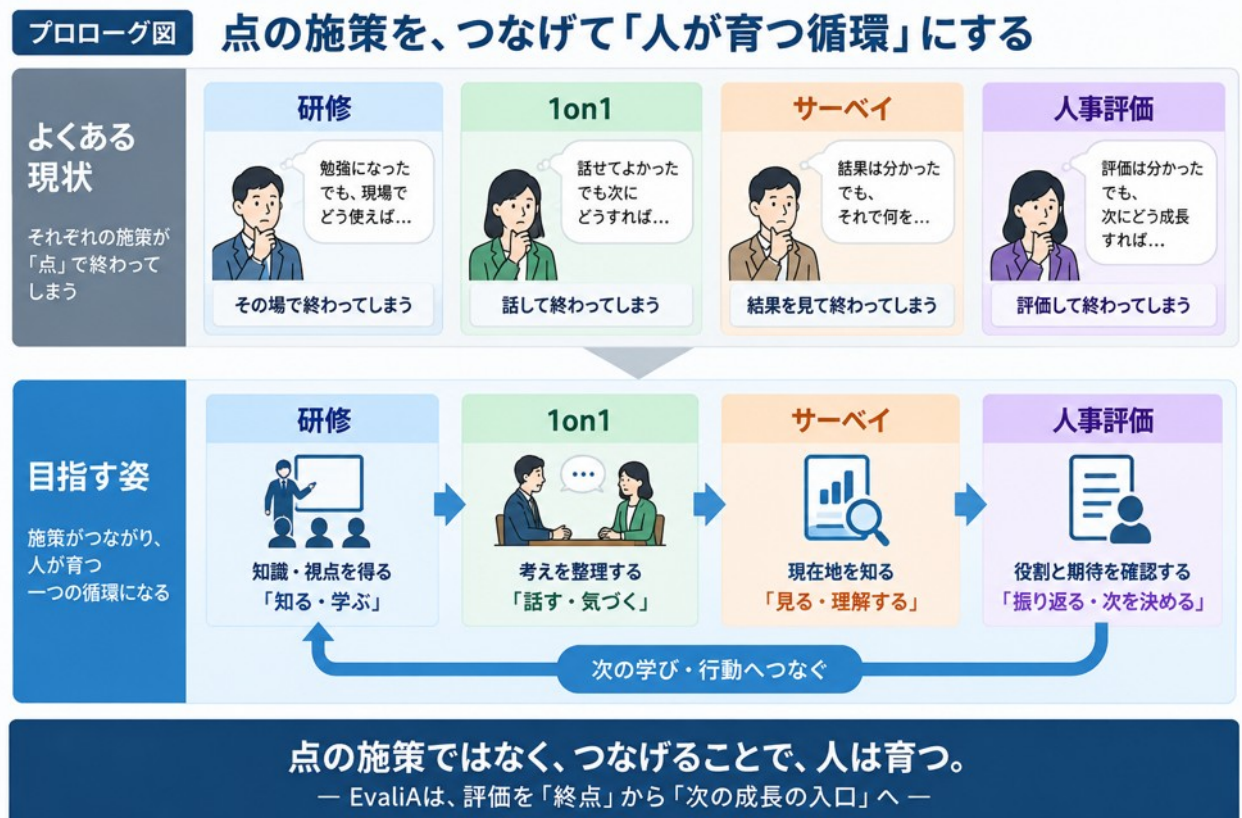
これらがつながって初めて、人材育成は「点」ではなく、閉じた学習循環になります。

センゲが示した「学習する組織」は、まさにこうした継続的な学習を組織全体で起こす考え方でした。しかし、概念を理解することと、一人ひとりについて日常業務の中で回し続けることは同じではありません。

人が多くなればなるほど、前回何を決めたのかを覚え、実際に何を試したのかを確認し、本人の状態を見て、成果と比較し、どこに原因があるのかを考え、次の行動へ戻すことは難しくなります。

ここで、AI の意味が出てきます。AI が人を育てるわけではありません。AI が、これまで分断されていた研修、対話、観測、評価の間をつなぎ、学習の循環を止めないようにする。

そして重要な判断点では、人間が意味づけし、選び、判断し、責任を引き受ける。これが、本書でいうヒューマン・イン・ザ・ループ（Human in the Loop）です。AIの処理・学習プロセスの中に人間が意図的に介入し、判断・評価・修正などを行う仕組みです。



その後、標準化、IT、AIと進む中で、ようやく私は、当時描いた図を「静止した図」ではなく、「動き続ける仕組み」として考えられるようになりました。

その具体的な仕組みについては、後の章で詳しく説明します。ここで伝えたいのは、一つだけです。

私は長い間、「どうすれば社員を育てられるのか」を探していました。

しかし、最後にたどり着いた問いは、「人をどう変えるか」ではなく、「人と会社がどう学び続けるか」でした。人が学ぶ。その人からチームが学ぶ。チームから会社が学ぶ。会社が変わる。変わった会社から、また人が学ぶ。

振り返ってみれば、私が本当に探していたのは、特別な教育手法ではありませんでした。

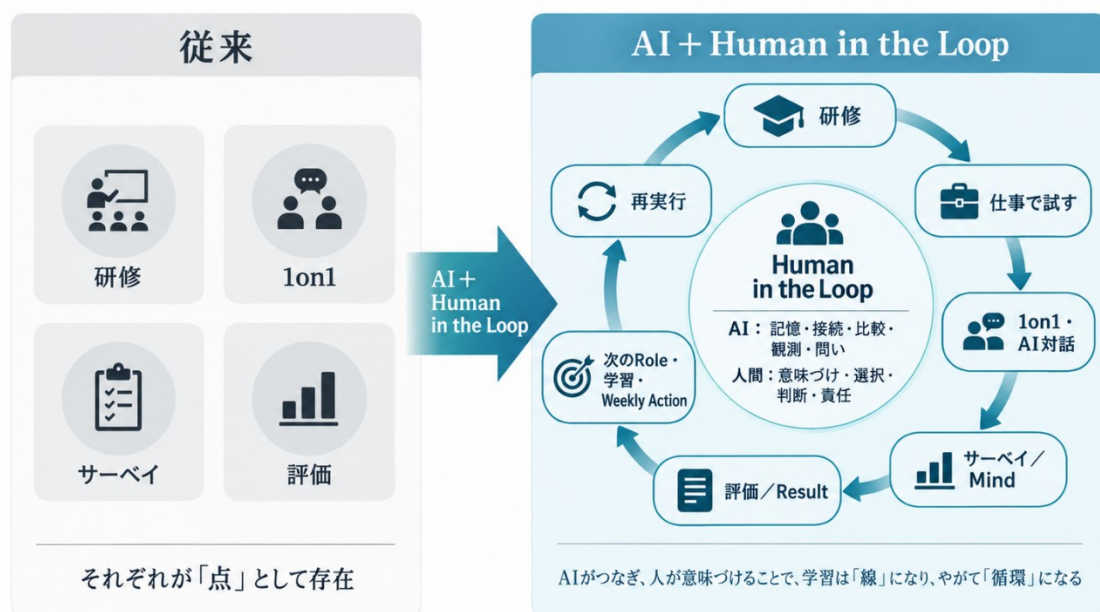
学習する組織(Learning Organization)だったのです。

では、経営学は、人と組織が学ぶという問題をどのように考えてきたのでしょうか。人を管理するという発想から始まった経営思想は、どのように人が自ら学ぶという考えへ進み、さらに組織そのものが学ぶという発想へたどり着いたのでしょうか。

経営学の歴史をたどりながら、なぜ現在の仕組みがこのような設計になったのかを考えていきます。

点だった人材施策が、AIで循環になる

—— 研修・1on1・サーベイ・評価を、閉じた学習ループへ ——



第一部 なぜ「学習する組織」を目指すのか

第1章 人を管理することから、学習する組織へ

——人間観の歴史と経営学 100 年の知見を、AI で一つの仕組みにする

プロローグで私は、長い間「どうすれば人を育てられるのか」を考え続け、最後には、「人を直接変える」のではなく、「人と会社が互いに学び続ける仕組み」をつくることの方が重要なのではないか、というところまでたどり着いた経緯を書きました。

しかし、これは私一人がぶつかった問題ではありません。

経営学の歴史を振り返ると、企業は百年以上にわたって、似た問いと格闘してきました。人は、なぜ期待したとおりに動かないのか。成果が出ないのは本人の努力不足なのか、それとも仕事の設計に問題があるのか。給与を上げれば人は動くのか。人間関係を良くすればよいのか。本人の考え方を変える必要があるのか。それとも、「人を変えよう」とする発想そのものを見直す必要があるのか。

経営学の歴史は、企業をどう管理するかという技術の歴史であると同時に、人間をどのような存在として見るのかという「人間観」の歴史でもあります。人をどう管理し、どう効率よく働かせるかという視点から始まり、人は合理性だけでは動かないという発見へ進み、さらに、人を外から動かすことより、人が自ら学ぶことへ、個人の学習からチームの学習へ、そして組織そのものが学ぶという考えへと視野は広がってきました。

ここで最初に明確にしておきたいことがあります。私は、これから紹介する理論を順番に勉強し、それらを組み合わせてAIアプリケーションである EvaliA をつくったわけではありません。実務の中で人材育成の成功と失敗を繰り返し、その結果として現在の仕組みが生まれました。後からその仕組みを経営学の歴史に照らしてみると、過去の研究者たちが考えてきた問題と、多くの接点が見えてきたのです。

したがって本章は、「この理論からこの機能をつくった」という発明の系譜を説明する章ではありません。EvaliA がなぜ現在のような設計になったのかを、経営学が積み上げてきた知見によって読み直す章です。

人は信じるべきか、管理すべきか

—— 孟子・荀子からマクレガーまで続く「人間観」の問題

人間をどのような存在として見るかという問いは、経営学が生まれるはるか以前から存在していました。中国思想には、孟子の性善説と荀子の性悪説があります。ただし、「孟子は人間を善人だと考え、荀子は悪人だと考えた」と単純化するのには正確ではありません。

孟子が重視したのは、人間には善へ向かう萌芽があり、それを育てることができるという考え方です。一方、荀子は、人間が自然に持つ欲望をそのままにすれば争いが起こり得るため、教育や礼、制度によって整える必要があると考えました。

二千年以上前から、私たちは同じ問いを持っていたことになります。人は信頼すべき存在なのか。それとも、制度によって行動を整える必要があるのか。

近代経営学の世界で、似た問いを「経営者が働く人をどう見るか」という形で整理したのが、ダグラス・マクレガーです。マクレガーは、管理する側が人間について持っている前提を、X 理論(Theory X)と Y 理論(Theory Y)として示しました。

X理論では、人は仕事を避ける傾向があり、成果を出すためには管理や統制が必要だという前提を置きます。一方のY理論では、適切な条件が整えば、人は自己統制し、責任を引き受け、自ら仕事に取り組むことができると考えます。

重要なのは、どちらが絶対に正しいかということではありません。経営者が人をどう見ているかによって、その会社の制度やマネジメントが変わり、その制度が社員の行動を変え、その行動を見た経営者が自分の人間観をさらに強めることです。

例えば、「社員は放っておけば自分では考えない」と上司が考えたとします。そこで仕事を細かく指示し、何をするにも承認を求めさせ、部下が自分で判断する余地を小さくする。すると部下は、本当に自分で考えなくなります。その姿を見た上司は、「やはり、社員は自分では考えない」と確信します。

人間観が制度をつくり、制度が行動をつくり、その行動が人間観を強化する。

ここには一つの循環があります。では私は、性善説と性悪説、X理論とY理論のどちらを選ぶのでしょうか。現在の私の答えは、どちらか一方ではありません。

制度は善意に依存させない。しかし、日々の運用では悪意を前提にしない。

人は、自ら成長しようとする可能性を持っています。一方で、人は忘れます。勘違いもします。疲れれば判断力も落ちます。自分自身を正確に認識できないこともあります。善意を持っていたても、ミスは起こります。

だから制度には、確認、承認、権限分離、異常検知などが必要です。しかし、問題が起きるたびに「この社員は怠けている」「責任感がない」と決めつけられれば、組織は監視と処罰を中心に動くようになります。人は失敗を隠し始めます。失敗が見えなくなれば、組織はそこから学ぶことができません。

この人間観は、EvaliAの基本設計にもつながっています。本人が何になりたいのかは本人が決める。次にどの行動を試すのかも、最後は本人が選ぶ。一方で、前回何を決めたのかを記憶し、同じ説明が何週間も繰り返されていないかを見たり、本人には見えにくい別の視点を返したりする部分は、AIが支援できます。

上司についても同じです。AIが最終評価を決めるのではありません。しかし、上司自身の思い込みだけで判断しないように、事実や本人の寄与を問い直すことはできます。

AIへ人間の判断を渡し切らない。しかし、人間の善意や記憶力だけでも依存しない。本書では、この考えをヒューマン・イン・ザ・ループ (Human in the Loop) と呼びます。これは、AIや自動化処理システムの処理プロセスに人間が意図的に関与し、確認、承認、判断、修正、責任の引き受けを行う仕組みです。このコンセプトを使って、「学習する組織」(Learning Organization) をどのように作るかを考えました。

テイラー —— 人を叱る前に、仕事を設計する

近代経営学の出発点の一つに、フレデリック・W・テイラーの科学的管理法 (Scientific Management) があります。テイラーというと、「人間を機械のように扱った」という印象を持つ人もいるでしょう。確かに、作業を細かく分け、時間を測り、効率的な方法を管理側が設計する考えには、そのような側面があります。

しかし、本書で注目したいのは別の点です。成果が出ないときに、「もっと一生懸命働け」と人へ要求するだけでなく、仕事そのものを観察し、分析し、より良い方法へ設計し直そうとしたことです。

作業が遅いなら、人を急がせる前に無駄な工程がないかを見る。ミスが多いなら、「もっと注意してください」と求めるだけでなく、仕事の手順や確認工程を見直す。

人を変える前に、仕事を変えられないか考える。この視点は、人材育成にもそのまま使えます。同じミスが繰り返されているなら、本人の注意力だけを問題にするのではなく、入力方法を変えられないか、確認工程を変えられないか、判断基準を明確にできないかを見る。

人の問題と、仕事の問題を分ける。

本書では、この考えを後にさらに広げていきます。本人を変えるべきなのか、上司を変えるべきなのか、チームを変えるべきなのか、それとも仕組みそのも

のを変えるべきなのか。成果が出ないときに、すべてを本人へ戻さないための重要な出発点です。

人間関係論 —— 個人だけを見ても、人の行動は分らない

仕事の設計だけでは、人の行動を十分には説明できません。

1920年代から1930年代に行われたホーソン研究(Hawthorne Studies)を契機として、職場の人間関係、集団規範、上司との関係などに関心が広がり、人間関係論(Human Relations)が発展しました。ホーソン研究の個々の結果については、その後さまざまな再検討があります。しかし、経営学の視野が「作業」だけから「人間関係や集団」へ広がったことには大きな意味があります。

例えば、ある社員が新しい意見を言わないとします。その人に主体性がないのでしょうか。

もしかすると、そのチームでは以前、新しい意見を出した人が上司から何度も否定されていたのかもしれませんが。「余計なことを言わない方が安全だ」という暗黙の規範ができている可能性もあります。

そうであれば、「もっと主体的になりなさい」と本人だけに言っても問題は解決しません。人の行動は、本人の内面だけではなく、その人が置かれているチームとの相互作用によって生まれるからです。

この視点が、本書で個人学習だけでなく、チーム学習(Team Learning)を重視する理由につながります。

後に詳しく説明しますが、私が概念化したCAGE適応モデル(CAGE Adaptation Model)も、個人を四種類の性格へ固定するためのものではありません。**成果(Challenger)、探究(Adventurer)、安全(Guardian)、関係(Empathizer)**という異なる行動機能が、チームの中で必要に応じて使われているかを見るためのモデルです。

一人の社員だけを変えて終わるのではなく、チームそのものを学習主体にする。人間関係論から本書が受け取るのは、この視点です。

マズロー —— 会社が望むことより先に、本人が望むことを聞く

仕事の設計や人間関係だけではなく、「そもそも人は何を求めているのか」という問いも研究されてきました。

アブラハム・マズローは、人間の欲求について、生理、安全、所属や愛、承認、自己実現という階層的な考え方を提示しました。現在では五段階のピラミッドとして広く知られていますが、すべての人が機械的に同じ順番で進む固定モデルとして理解するのは適切ではありません。

本書で重要なのは、人間が「給与を与えれば働く存在」だけではないという人間観です。人は生活するために働きます。しかし、それだけではありません。仲間に認められたい。自分の能力を生かしたい。成長したい。仕事に意味を感じたい。将来こうなりたいという姿を持ちたい。

ここから、人材育成にとって一つの大切な問いが生まれます。会社が社員へ「この目標を達成してください」と伝える前に、本人自身が何を望んでいるのかを聞く必要があるのではないか。

もちろん、「本人のありたい姿から始める」という EvaliA の設計が、マズローの理論から直接導かれたわけではありません。しかし、人間を会社から与えられる報酬だけで動く存在として見ないという点で、共通する人間観があります。

だから EvaliA では、会社から与えられた重要業績評価指標(Key Performance Indicator)だけから始めません。

最初に本人へ、「あなたは、何になりたいのですか」と問いかけます。

この問いが、後に述べる自己マスタリー(Personal Mastery)の出発点になります。

ハーズバーグ —— 「不満がない」と「成長している」は違う

フレデリック・ハーズバーグは、仕事への不満に関係しやすい要因と、仕事そのものへの積極的な満足や動機づけに関係する要因を区別しました。一般に、給与、会社方針、監督、勤務環境などは「衛生要因」(Hygiene Factors)、達成、承認、仕事そのもの、責任、成長などは「動機づけ要因」(Motivators)と呼ばれます。

この理論を、人間心理をすべて説明する普遍的な法則として扱う必要はありません。しかし、人の状態を考える一つのレンズとして、重要な問いを与えてくれます。

不満がないことと、人が成長していることは同じなのか。

給与にも大きな不満がない。上司との関係にも問題はない。職場環境も悪くない。しかし、本人が何年も新しい仕事へ挑戦せず、成長実感もない。その人を「良い状態」とだけ判断してよいのでしょうか。

反対に、難しい仕事へ挑戦しているため一時的に苦しさがあっても、本人が強い達成感や成長実感を持っていることもあります。

この違いは、後に本書で心理状態(Mind)を扱うときに重要になります。

人の状態を一つの点数で決めない。何に満足しているのか。何が不満なのか。何に意味を感じ、何を成長と感じているのか。その状態だけで結論を出さず、「なぜそうなっているのか」を考える。

ハーズバーグから本書が受け取るのは、この視点です。

アージリス —— 行動だけでなく、前提そのものを問い直す

—— 単一循環学習(Single-loop Learning)から二重循環学習

(Double-loop Learning)へ

クリス・アージリスは、組織学習を考えるうえで重要な区別をしました。単一循環学習(Single-loop Learning)と二重循環学習(Double-loop Learning)です。

単一循環学習では、基本的な目標や前提を維持したまま行動を修正します。例えば営業目標を達成できなかったときに、「来週は訪問件数を増やそう」と行動を変える。これは必要な学習です。

しかし、二重循環学習では、さらに深く問い直します。

そもそも、この顧客へ訪問することが正しいのか。対象顧客そのものが間違っていないか。「訪問件数を増やせば成果が上がる」という前提自体は正しいのか。

つまり、行動だけではなく、**その行動を生み出している前提そのものを学習対象にする**のです。

例えば、ある社員が三週間続けて「時間がなくてできませんでした」と話していたとします。「では来週は時間を確保しましょう」だけなら、行動修正で終わります。

しかし、過去の対話を比較すると、毎週、緊急業務を優先して、重要だが緊急ではない行動が後回しになっているかもしれません。

そこでAIは、「三週間、同じパターンが続いています。これは時間そのものの不足だけではなく、優先順位のつけ方が関係している可能性はありませんか」と問い返すことができます。

さらに、「そもそも、現在設定している目標は、今のRoleに適しているのでしょうか」と前提そのものへ問いを広げることができます。

ここで重要なのは、AIが原因を断定することではありません。本人がこれまで見ていなかった問いを返すことです。AIを「答えを教える教師」ではなく、**自分の前提を映す鏡**として使う。

後の章で扱うメンタルモデル(Mental Models)やGapの考え方は、ここへつながっていきます。

センゲ —— 最上位に置くのは「学習する組織」である

—— 五つのディシプリン(Five Disciplines)を設計の定規にする

ここまで紹介してきた理論は、それぞれ人や組織の一側面を照らしています。しかし、本書では、それらを最終的に一つの上位概念の中へ置きます。

それが、ピーター・センゲの学習する組織(Learning Organization)です。私たちが目指しているのは、エンゲージメントが高い会社だけでも、AIを導入した会社でも、研修が充実した会社でもありません。

人が学び、チームが学び、組織が学び、その結果としてまた人が学ぶ会社です。そして、その組織が本当に学習する能力を持っているのかを見る定規として、センゲの五つのディシプリン(Five Disciplines)を使います。

第一は、自己マスタリー(Personal Mastery)です。自分が何を実現したいのかという方向を持ち、現在の現実との間にある差を認識しながら、学び続けます。本書では、本人へ「何になりたいのか」と聞くところから、この学習を始めます。

第二は、メンタルモデル(Mental Models)です。人は現実そのものを見てい
るのではなく、過去の経験や価値観からつくられた前提を通して現実を見ています。「私は十分に部下へ任せている」「私は顧客志向である」と本人が考えていても、実際の行動は異なるかもしれません。自分自身のものの見方も、学習対象になります。

第三は、共有ビジョン(Shared Vision)です。本書では、その実務への第一歩として、本人が目指す未来と、会社が目指す未来との接点を見つけます。会社が理念を一方的に伝えて終わるのではなく、「自分が目指すものと会社が目指すものは、どこでつながっているのか」を本人自身が考えます。

第四は、チーム学習(Team Learning)です。個人が学んだだけでは、組織にはなりません。異なる視点を持つ人たちが、自分たちの働き方そのものを学習対象にし、役割や協力関係を変えながら学ぶ必要があります。

第五が、システム思考(Systems Thinking)です。社員が成果を出せないとき、その瞬間に「本人の能力不足だ」と結論づけない。本人、上司、顧客、チーム、制度、時間軸から同じ現象を見る。他の社員にも同じことが起きていないかを見る。役割や目標設定そのものが正しいのかも見る。

一つの結果を、複数の相互作用の中で捉える。これがシステム思考です。

五つのディシプリンは、五つの研修科目ではない

ここで重要なのは、五つのディシプリンを五つの独立した研修科目として考えないことです。自己マスタリー研修、メンタルモデル研修、共有ビジョン研修、チーム学習研修、システム思考研修を順番に行えば学習する組織ができる、というものではありません。

本人が自分の未来を考える。現在地を見る。自分がどのような前提で現在地を見ているのかを問い直す。会社が目指す方向との接点を探す。チームの中で役割(Role)を担う。実際に行動する。その結果を振り返る。本人だけでは解決できない問題なら、チームや仕組みへ視点を広げる。

この一連の仕事の中で、複数のディシプリンが同時に働きます。したがって、EvaliAの中に「自己マスタリー機能」「メンタルモデル機能」といった五つの独立した箱があるわけではありません。

五つのディシプリンは EvaliA の五機能ではない。EvaliA の設計が、本当に学習する組織へ向かっているかを検証する五つの定規なのです。

野中郁次郎 —— 個人の気づきを、組織の知識へ変える

個人が学んだだけでは、会社が学んだことにはなりません。ある熟練者が、「この顧客の場合は、最初にここを確認した方がよい」と経験的に分かっている、その知識が本人の頭の中にしかなければ、その人がいなくなったときに会社から消えてしまいます。

ここで重要な補助線になるのが、野中郁次郎らの知識創造論です。知識創造の四過程モデル(SECI Model)では、共同化(Socialization)、表出化(Externalization)、連結化(Combination)、内面化/Internalization)という循環を通じて、暗黙知と形式知が相互に変換されていきます。

経験する。経験の中にあつた知識を言葉にする。複数の知識を結びつける。それを別の人実践し、また新しい実践知を生み出す。

本書でも、個人の経験を本人の中だけで終わらせないことを重視します。本人が仕事の中で気づいたことを言語化し、複数人に共通するパターンがあれば組織側へ返す。そして、仕事の標準や仕組みを変える必要があるなら、改善へつなげる。

そのための出口として、本書では共有(Share)→確認(Check)→調整(Adjust)→実行(Do)を回す SCAD を位置づけます。これも私が概念化したボトムアップの改善手法です。詳しくは、拙著、「やっぱり仕組み化」(TCG 出版 2024 年) 参照。

EvaliA が個人の気づきを生み、SCAD がその気づきを組織学習へつなぐ。個人の学習を組織の知識へ変えるためには、この出口が必要です。

カーン —— エンゲージメントは「状態」を見る

近年、人事の世界ではエンゲージメント(Engagement)という言葉が広く使われています。その概念を考えるうえで重要な研究者の一人が、ウィリアム・カーンです。

カーンは、人が仕事上の役割にどの程度自分自身を投入できるかという観点から、意味、安全、利用可能性などの心理的条件を考えました。本書で重要なのは、エンゲージメントを単なる「会社への満足度」として扱わないことです。

人と仕事、人と Role、人と職場との関係から生まれる状態を見る。ただし、状態が分かっただけでは原因は分かりません。

エンゲージメントが低い。その原因は上司との関係かもしれません。Role が本人の未来と合っていないのかもしれない。仕事量が過大なのかもしれない。成長を感じられていないのかもしれません。

だから、エンゲージメントを高めること自体を最終目的にはしません。エンゲージメントは「状態を見る」。ハーズバーグは「なぜその状態なのかを考える」。本書では、その先の「では、どこを変えるのか」まで進みます。

ジェームズ・リーズン —— 人を責めるのか、仕組みを直すのか

人が失敗したとき、私たちはつい、「誰が間違えたのか」を探します。

しかし、もう一つ重要な問いがあります。「なぜ、その間違いが失敗として通過できたのか」。

人間のエラーや組織事故を研究したジェームズ・リーズン(James Reason)は、人の失敗を見る考え方として、「人的アプローチ」(Person Approach)と「システムアプローチ」(System Approach)を対比しました。

人的アプローチでは、注意不足だった、忘れた、判断を誤った、と個人側へ原因を求めます。一方、システムアプローチでは、人間が間違えることを前提として、なぜ手順で防げなかったのか、なぜ確認工程で止まらなかったのか、なぜ複数の防御を通過して結果まで到達したのか、と仕組みや組織の側まで原因探索を広げます。

本書では、この視点の転換を、「ヒューマンエラー」(Human Error)から「システムエラー」(System Error)へ原因探索を広げることと表現します。ただし、「システムエラー」はここで本書が用いる整理であって、リーゼンの用語そのものという意味ではありません。

また、すべての失敗について本人に責任がないという意味でもありません。**故意の違反は人の責任として扱う。故意でないミスは、仕組みが学ぶ材料として扱う。**ここを分ける必要があります。

日本の製造現場で発展したポカヨケ(Poka-Yoke)も、人へ「間違えないように注意してください」と求めるだけではなく、そもそも間違えにくくする、あるいは間違えた瞬間に分かるようにする、という仕事の設計です。

本書で重視するのは、個々のポカヨケそのものより、その背後にある考え方です。人の失敗を、本人だけの問題として終わらせず、仕組みが学ぶきっかけへ変える。

人を叱って終わる組織では、失敗は消費される。失敗から仕組みを変える組織では、失敗が学習資産になる。後の章では、この循環を失敗防止循環(Foolproof Loop)として具体化します。

ナッジ ——心を直接変える前に、小さな行動を変える

私自身の人材育成を見直すうえで、もう一つ重要な考え方がナッジ(NUDGE)です。

ナッジでは、人の選択肢を奪って強制するのではなく、本人の選択を残したまま、ある行動を選びやすい環境をつくります。これは、以前の私の人材育成とはかなり違います。

以前は、「もっと主体性を持ちなさい」「顧客志向になりなさい」「経営者マインドを持ちなさい」と、本人の内面へ直接働きかけようとしていました。

しかし現在は、例えば、「次に上司へ相談する前に、自分の案を一つ考えるとしたら何をしますか」と、小さな行動へ落としします。

「顧客志向になりなさい」ではなく、「次の顧客対応で、回答する前に一つだけ相手の目的を確認する」としたら、何を聞きますか」と問いかける。

最初から価値観を変える必要はありません。まず一つ行動を変える。行動が変われば経験が変わる。経験を振り返れば、本人自身が何かに気づくことがあります。その気づきが、やがてものの見方を変えることもあります。

つまり、人の学習には、「**考え方→行動**」だけではなく、「**行動→経験→内省→考え方**」という方向もあります。

本書では、ナッジの考え方を AI コーチングの設計へ取り入れます。ただし、ナッジそのものと、後に述べる週次行動(Weekly Action)は同じものではありません。ナッジは本人が行動を選びやすくする設計原則であり、週次行動は、その中で本人が選ぶ具体的な学習実験です。

振り返ってみると、ゴールはすでに見えていた

ここまで経営学の流れを見てきて、以前の著書を読み返すと、興味深い記述があります。私は当時、行動することによって経験が増え、経験によって質的な変化が生まれ、それが成長につながり、成長が次の行動を促すという循環を描いていました。そして、その循環が組織的に起きている企業であれば、「学習する組織」と言えるだろうとも書いていました。

さらに、経験、成長、モチベーションといった内面は外から直接見ることができないため、外に表れた行動やその変化を見る仕組みが必要だとも考えていました。

二十年ほど前の私は、ゴールのかなり近くまで来ていたのだと思います。

人は行動から学ぶ。その循環が組織全体で起きれば、学習する組織になる。内面は直接見えないから、外に表れた変化を見る。

しかし、一番難しい問題が残っていました。**これを、一人ひとりについて、どうやって継続的に回すのかです**。理念はありました。必要な理論も、世の中にはかなり揃っていました。しかし、それらを日常業務の中で動かし続ける運転能力がありませんでした。

AI 変えたのは、「答え」よりも「継続」である

AI については、「人間より賢い答えを出せるか」「どの仕事を自動化できるか」という議論が多くあります。もちろん、それも重要です。

しかし、学習する組織という観点から見たとき、私は人工知能(AI)の価値を別のところに感じています。**AI が大きく変えたのは、「答え」よりも「継続」です。**

例えば社員が 300 人いる会社を考えてみます。一人ひとりが何を目指しているのかを確認する。前回何を決めたのかを踏まえ、今週何を実行したのかを聞く。できなかったなら、その理由を考える。本人だけではなく、上司、顧客、チーム、組織、時間軸と視点を変える。何週間も同じところで止まっていれば、そのパターンを本人へ返す。複数人に同じことが起きていれば、個人ではなく構造を疑う。そして、しばらく後に本当に変化したのかをもう一度見る。

これを管理職だけで、全社員について、毎週高い品質で継続することは容易ではありません。AI によって、その継続運用が現実的になってきました。

ただし、人工知能(AI)に人生を決めさせるわけではありません。本人が何を目指すかを選ぶ。上司が誰にどの Role を任せるのかを判断する。本人が次に何を試すかを決める。上司が成果の意味を判断する。チームが何を学ぶかを考える。会社がどの仕組みを変えるかを決める。

人工知能(AI)は、その間をつなぎ、観測を続け、比較を続け、問いを止めない。

学習する組織に足りなかったのは、必ずしも新しい理論ではなかったのかもしれない。**足りなかったのは、学習を継続的に運転する能力でした。**

AI は、その運転コストを大きく変えたのです。本書のタイトルにある「簡単につくれた」という言葉も、学習する組織をつくること自体が簡単になったという意味ではありません。これまで一部の優れた上司の力量や膨大な時間に依存していた問い、振り返り、比較を、一人ひとりについて繰り返し回せる可能性が高まった。

その意味での「簡単」です。

経営学 100 年の知見を、一つの学習循環へ

ここまで見てきた理論は、別々の時代に生まれた、別々の理論です。

しかし私には、それぞれが「人と組織が学び続けるために必要な視点」を、一つずつ見つけてきた歴史のように見えます。孟子・荀子からマクレガーへと続く人間観からは、**信頼と制度を両立すること**。

テイラーからは、**人を責める前に仕事を設計すること**。

人間関係論からは、**個人だけを見ても人の行動は理解できないこと**。

マズローからは、**会社の要求だけでなく、本人が何を実現したいのかを見ること**。

ハーズバーグからは、**不満がないことと成長していることを分けること**。

アージリスからは、**行動だけでなく、行動を生み出している前提まで問い直すこと**。

センゲからは、これらを個人だけの学習で終わらず、**学習する組織 (Learning Organization)**という上位概念で捉えること。

野中郁次郎らの知識創造論からは、**個人の経験を組織知へ変えること**。

カーンをはじめとするエンゲージメント研究からは、**人と仕事との関係から生まれる状態を見ること**。

ジェームズ・リーズンからは、**失敗を本人だけの問題にせず、仕組み側まで原因探索を広げること**。

そしてナッジからは、**価値観を直接変えようとせず、小さな行動から学習を始めること**。



これらの理論を、EvaliA の機能へ一対一に割り当てたわけではありません。それぞれが示した問題を、一つの学習循環の中で扱おうとしているのです。

本人が何をを目指すのかを考える。現在地を見る。自分がどのような前提で現在地を見ているのかを問い直す。会社が目指す方向との接点を探す。チームの中で Role を持つ。小さな行動を試す。行動、本人の受け止め、成果を異なる角度から見る。

その間に違いがあれば、本人の問題なのか、上司なのか、チームなのか、仕組みなのかを考える。個人の気づきを必要に応じて組織へ戻す。そして、変わった組織の中で、また本人が学ぶ。

ここで大切なのは、AI によって人を細かく監視することではありません。人と組織が学ぶために必要な変化を、継続的に観測することです。

本人が自分を振り返るための空間と、組織が学ぶために使う情報は分ける必要があります。AI が見られるからといって、すべてを上司が見てよいわけではありません。

監視によって「正しい答え」を言わせるのではなく、観測した変化を本人や組織へ返し、次の問いを生み出す。私は、この違いを非常に重要だと考えています。

EvaliA は、単なる AI コーチング・システムではありません。経営学が長い時間をかけて積み上げてきた「人と組織が学ぶための知見」を、日常業務の中で一つの循環として動かそうとする試みです。

では、その循環はどこから始まり、どのように個人からチームへ、さらに組織へ広がっていくのでしょうか。その出発点として、私は二十年ほど前に描いた一枚の図へ戻ることになりました。

次章では、その「個人目標」と「組織目標」という二つのベクトルを出発点に、学習する組織(Learning Organization)を日常業務の中で運転する全体構造を見ていきます。

第2章 20年前の「二つのベクトル」をAIで回す

——学習する組織(Learning Organization)を、運転可能な仕組みに変える

前章では、経営学が長い時間をかけて、人を管理するという発想から、人が自ら学び、チームが学び、さらに組織そのものが学ぶという考えへ進んできた流れを見てきました。本書では、その最上位に学習する組織(Learning Organization)を置きます。

そして、その組織が本当に学習する能力を持っているかを見る定規として、自己マスタリー(Personal Mastery)、メンタルモデル(Mental Models)、共有ビジョン(Shared Vision)、チーム学習(Team Learning)、システム思考(Systems Thinking)という五つのディシプリン(Five Disciplines)を使います。

しかし、ここからが実務の問題です。学習する組織という考え方を知ることと、それを実際の会社の中で動かし続けることは同じではありません。自己マスタリーを高めよう、メンタルモデルを問い直そう、共有ビジョンをつくろう、チームで学ぼう、システム思考を身につけよう、と言うことはできます。難しいのは、それを一人ひとりについて、日常の仕事の中で繰り返すことです。

私は、この問題を考えたとき、二十年ほど前に自分が描いた一枚の図へ戻ることになりました。そこに描かれていたのが、「個人目標」と「組織目標」という二つのベクトルです。

二十年前に描いた「個人目標」と「組織目標」

前掲書で私は、「個人目標」と「組織目標」という二つの方向によって、会社の組織力を考えていました。ここでいう個人目標とは、会社の売上目標や予算を個人へ配分したものではありません。本人が何をやりたいのか、何を目指

しているのか、どのような人になりたいのか。そうした本人自身の方向を意味していました。

一方の組織目標は、経営理念、ビジョン、方針、戦略、予算など、会社がどこへ向かおうとしているのかを示すものです。会社がどれほど立派なビジョンを掲げていても、そこで働く本人が、その方向に自分自身の意味を見いだせなければ、大きな力にはなりません。

反対に、一人ひとりが自分のやりたいことだけを追いかけ、会社が目指す方向とまったく接点がなければ、組織として一つの力になることも難しくなります。

よって私は、個人の方向と組織の方向が重なることが重要だと考えていました。この基本的な考え方は、今でも変わっていません。

ただし、二十年前の図には、一つ大きなものが欠けていました。それは**時間**です。

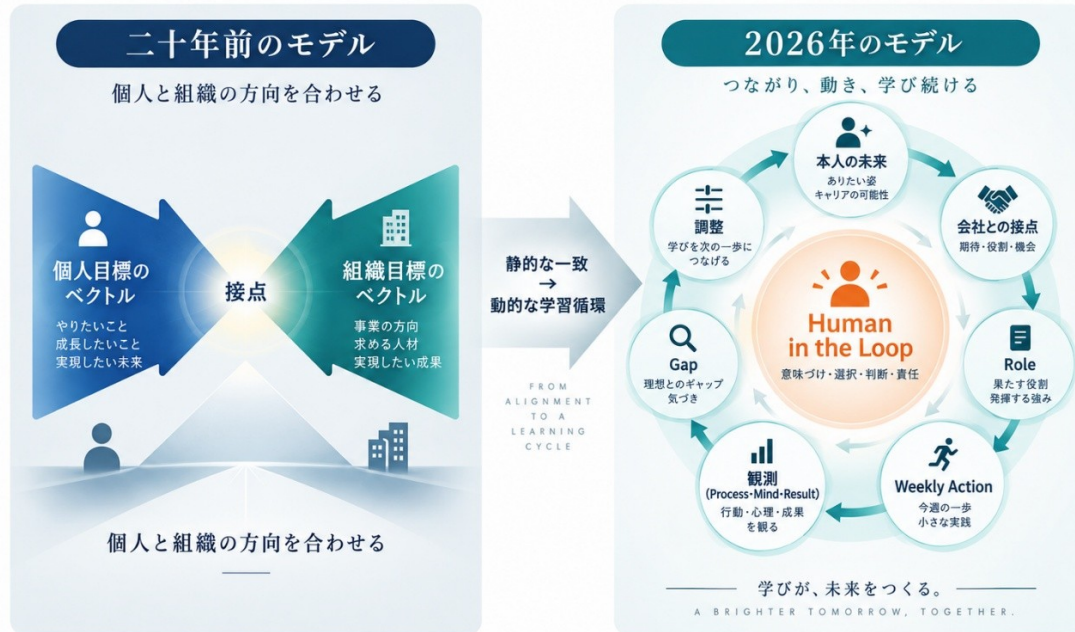
人は変わります。会社も変わります。市場環境も、上司も、チームの構成も変わります。当然、本人に求められる役割(Role)も変わります。一度、個人と会社の方向が重なったとしても、一年後にも同じように重なっているとは限りません。

したがって、本当に必要なのは、一度二つのベクトルを一致させることではありません。**ずれを観測し、そのずれから学び、何度でも調整し直すこと**です。二十年前の静的な「二つのベクトル」を、動き続ける学習循環へ変える。ここが現在の EvaliA の出発点です。

図2-1

二十年前の二つのベクトルから、2026年の学習循環へ

FROM TWO VECTORS TO A LEARNING CYCLE



一本目のベクトル —— 本人はどこへ行きたいのか

二十年前に私が「個人目標」と呼んでいたものを、現在の視点から読み直すと、自己マスタリー(Personal Mastery)の出発点と重なります。自己マスタリーとは、単に知識や技術を増やすことではありません。自分は何を実現したいのか、どのような人になりたいのか、そして現在の自分はどこにいるのか。この「ありたい姿」と「現在地」の差を認識しながら、学び続けることです。

EvaliA では、会社が本人へ与えた重要業績評価指標(Key Performance Indicator)だけから始めません。最初に本人へ、「あなたは、何になりたいのですか」と問いかけます。

例えば、「将来は管理職になりたい」と答える人がいたとします。しかし、それだけではまだ十分ではありません。なぜ管理職になりたいのでしょうか。人を育てたいのでしょうか。もっと大きな仕事をしたいのでしょうか。自分の専門性を生かして、より大きな価値を生み出したいのでしょうか。それとも、年齢的に次は管理職になるものだと思っているだけなのでしょうか。

AIは、正しい人生目標を教えるわけではありません。問いを重ねることによって、本人が自分自身の方向を言葉にすることを支援します。本人がどこへ行きたいのか。そこから学習が始まります。

ありたい姿だけでは足りない —— 「現在地」が必要になる

しかし、ありたい姿が分かっただけでは、成長は始まりません。もう一つ必要なのが、現在地です。そして、人材育成の難しさはここに 있습니다。人は、自分自身を必ずしも正確には見ていないからです。

「私は部下の話をよく聞いている」と思っている、実際には、部下が話し始めるとすぐに自分の答えを話しているかもしれません。「私は顧客志向だ」と思っている、自分が説明したいことばかり話しているかもしれません。「私は部下へ任せています」と言いながら、重要な判断はすべて自分で行っているかもしれません。

ここで関係してくるのが、メンタルモデル(Mental Models)です。自己マスタリー(Personal Mastery)が「私はどこへ行きたいのか。その未来に対して現在どこにいるのか」を扱うとすれば、メンタルモデル(Mental Models)は、「私は、その現在地をどのようなレンズを通して見ているのか」を問い直します。

本人から見る。上司から見る。部下から見る。顧客から見る。さらに時間軸を変えて見る。一つの見方だけを正しいものとしない。ここでは、自己マスタリーとメンタルモデルだけではなく、すでにシステム思考(Systems Thinking)も働き始めています。

つまり、五つのディシプリン(Five Disciplines)は、一つずつ順番に動くものではありません。一つの仕事、一つの対話の中で、複数の同時に働くのです。

なお、人の「現在地」を何によって見るのかという問題は非常に重要です。本書では、その一つの実務的な座標として、私が以前から使ってきた「心」と「技」の人財マトリックスを読み直します。この点は次章で詳しく考えます。

二本目のベクトル —— 会社はどこへ行くのか

本人がどこへ行きたいのかが見えても、それだけでは会社という組織にはなりません。会社にも方向があります。経営理念、ビジョン、戦略、部門目標があり、その中で一人ひとりに期待する役割(Role)があります。ここで、二本目のベクトルである「組織目標」が出てきます。

ただし、会社にビジョンが存在するだけでは、共有ビジョン(Shared Vision)にはなりません。会社が立派なビジョンをつくり、それを社員へ説明し、社員が内容を覚える。それだけで共有ビジョンが成立するのであれば、組織づくりはもっと簡単だったはずです。

重要なのは、本人が「自分が目指していることと、この会社が目指していることには、ここに接点がある」と、自分自身の意味として理解できることです。

例えば、本人が「将来、海外拠点の経営を任されたい」と考えていて、会社も海外展開を強化しようとしているのであれば、現在の海外業務は単なる会社都合の仕事ではなく、本人が目指す将来へ近づくための経験にもなります。

もちろん、いつも接点が見つかるとは限りません。その場合、無理に一致させる必要はありません。「現在はずれている」こと自体を見えるようにすることにも意味があります。

学習する組織(Learning Organization)とは、全員が最初から同じ方向を向いている会社ではありません。違いに気づき、対話し、必要なら個人も会社も方向を問い直せる会社です。

共有ビジョン(Shared Vision)を、チームの仕事へ落とす

個人と会社の方向に接点が見つかったとしても、まだ日常の仕事にはなっていません。次に必要になるのが、「では、あなたは、このチームの中で何を担うのか」という問いです。ここで役割(Role)が入ります。

役割は、単なる肩書ではありません。「営業担当」「マネージャー」「経理担当」といった職種名だけでもありません。本書でいう Role とは、「**このチームが、現在の目標を実現するために、あなたに何を担ってほしいのか**」という期待です。

しかも Role は、本人の能力だけを見て決めるものではありません。上司は、「この人は何が得意か」だけでなく、「今のチームには、どの機能が足りていないのか」を見る必要があります。

そのための一つの実務モデルとして、本書では CAGE 適応モデル(CAGE Adaptation Model)を使います。CAGE では、成果(Challenger)、探究(Adventurer)、安全(Guardian)、関係(Empathizer)という四つの行動機能を考えます。成果へ向かって決断し動かす。新しい可能性を探り、問いを立てる。リスクを確認し品質や安定を守る。人と人をつなぎ協力関係をつくる。この四つです。

ここで重要なのは、CAGE を「四種類の人間」に分類するために使わないことです。今、このチームには何が必要なのか。その中で本人には、どの機能を主に担ってほしいのか。この順番で Role を考えます。

したがって、全体の流れは、**共有ビジョン(Shared Vision)→チーム目標→必要な CAGE 機能→役割(Role)**となります。ここで初めて、個人の未来と会社の未来が、チームの日常業務へつながり始めます。CAGE 適応モデルの詳しい考え方は、第 5 章で扱います。

Role を、一週間で試せる行動まで小さくする

Role が決まっても、そのままではまだ抽象的です。例えば上司が、「あなたには、もっと探究してほしい」と言ったとしても、本人からすれば、「具体的には何をすればよいのでしょうか」となります。そこで Role を、実際に試すことのできる行動へ落としていきます。

例えば、「今週、顧客一人に、これまで聞いたことのない質問を一つするとしたら、何を聞きますか」と問いかける。Role が部下育成であれば、「今週、一つの判断について、自分が答えを言う前に、部下の案を先に聞くとしたら、どの仕事で試しますか」と聞くこともできます。

重要なのは、AIが本人へ命令することではありません。問いや選択肢を返し、最後は本人自身に選んでもらいます。ここにナッジ(NUDGE)の考え方があります。

つまり、大きな流れは、**役割(Role)→期待行動→本人の選択→週次行動(Weekly Action)**です。会社の大きなビジョンが、最後には「今週、自分は何を一つ試すのか」という大きさまで下りてきます。ナッジと週次行動については、第6章で詳しく扱います。

学習の中心にあるのは、毎週の短い循環

一週間後、再びAIとの対話が始まります。前回、自分は何を試すと決めたのか。実際にやってみたのか。何が起きたのか。何を学んだのか。できなかったなら、何が止めたのか。別の見方はできないのか。本人だけの問題なのか、それとも上司やチーム、仕事の仕組みにも原因があるのか。こうした問いを通じて、経験を学習へ変えていきます。

EvaliAの中心にある短い循環は、**週次行動(Weekly Action)→実行→週次AIコーチング→振り返り→次の週次行動(Weekly Action)**です。この循環を毎週繰り返します。

ただし、重要なのは、この循環が本人の中だけで終わらないことです。本人自身で変えられる問題なら、本人の次の行動へ戻します。チームとして考えるべき問題なら、チーム学習(Team Learning)へ上げます。会社の仕組みを変えるべき問題なら、組織学習へ上げます。

つまり、週次AIコーチングは単なる個人コーチングではありません。**個人学習、チーム学習、組織学習へ分岐する起点**でもあるのです。

個人からチームへ、チームから組織へ

本人自身で変えられる問題なら、次の週次行動(Weekly Action)へ戻します。行動し、経験し、振り返り、また行動する。この短い循環が個人の学習です。

一方、複数人の行動を見たとき、「成果へ向かう行動は強いが、安全確認が少ない」「新しい提案は多いが、実行責任を誰も担っていない」「人間関係は

良いが、必要な対立を避けている」といった共通するパターンが見えることがあります。この場合、それを一人の能力不足として扱うのではなく、チーム全体の働き方として見る必要があります。

AI がパターンを整理し、上司やチームへ問いを返し、Role や協力関係、チームとしての行動を調整する。ここでチーム学習(Team Learning)が始まります。

さらに、同じ入力ミスが何人にも起きている、同じ情報を何人も重複して入力している、判断基準が人によって違う、何人もの社員が同じ承認待ちで止まっている、といった問題であれば、「本人が次から気をつける」だけでは足りません。仕事の仕組みそのものを変える必要があります。

その場合には SCAD へつなぎます。SCAD は、共有(Share)→確認(Check)→調整(Adjust)→実行(Do)を回す改善循環です。ここでの位置づけが重要です。SCAD は、人材成長循環(Human Development Loop)と並んで動く別の大きな仕組みではありません。**個人やチームの学習を進めた結果、「仕組みそのものを変える必要がある」と分かったときに起動する組織改善のサブループ**です。

EvaliA が個人やチームの気づきを生み、SCAD が必要な気づきを仕組みの変更へつなげる。そして、変わった仕組みの中でまた人が行動し、その結果を再び観測する。ここで、人の学習と組織の学習が一つにつながります。

AI は、質問するだけではない

週次 AI コーチングの価値は、毎週質問できることだけではありません。もう一つ重要なのが、過去と現在を比較できることです。

本人が「私は部下へ十分任せています」と話している一方、数週間の対話を見ると、重要な判断は毎回本人が行っているかもしれません。あるいは、「今週も時間がありませんでした」という説明が三週間続いているかもしれません。

そのとき AI は、「以前の発言と現在の行動の間に違いがあるように見えます」「三週間、同じところで止まっています」と、本人だけでは気づきにくいパターンを返すことができます。本書では、この働きを内省支援フィードバック(Reflective Feedback)と呼びます。

AIが「あなたはこういう人です」と判定するのではありません。「こういうパターンが見えます。あなたはどう考えますか」と、鏡のように返す。これによって、メンタルモデル(Mental Models)への気づきが生まれる可能性があります。

また、同じパターンが複数人に見つかれば、個人の問題からチームや組織の問題へ視点を広げることできます。AIが答えを決めるのではなく、問いの範囲を広げるのです。

学習を三つの窓から見る

ただし、毎週の行動だけを見ても、人の状態をすべて理解できるわけではありません。そこで本書では、人の学習を三つの異なる窓から見ます。

第一が行動過程(Process)です。本人は何を試したのか。何を実行したのか。前進したのか、停滞したのか。何が阻害要因になったのか。

第二が心理状態(Mind)です。本人は、現在の仕事をどのように受け止めているのか。何に満足し、何に不満を感じ、何に意味や負担、成長実感を得ているのか。

第三が成果(Result)です。会社が期待した成果は実際に出たのか。本人は何に寄与したのか。未達にはどのような理由があったのか。

ここで最も重要なのは、三つを最初から一つの点数に混ぜないことです。同じ本人の同じ仕事を、異なる窓から見た後で、その間にあるずれ(Gap)を見ます。

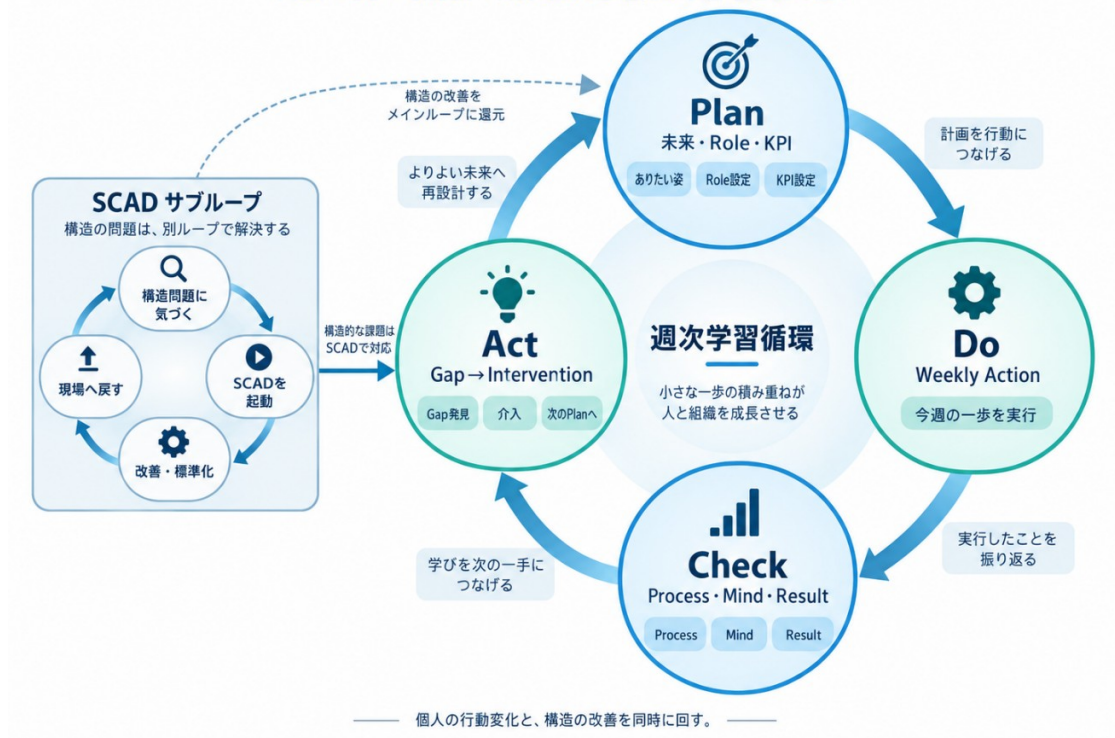
行動は前進しているのに、本人は強い負担を感じている。本人は順調だと感じているのに、成果が出ていない。成果は出ているが、その過程には再現性がない。こうした違い自体が、次に何を見るべきかを教えてくれます。

ただし、Gapが見えたからといって、原因まで分かったわけではありません。**Gapは原因ではなく、「次に、どこを詳しく見るべきか」を示す探索地点です。**

三経路観測とGapについては、第7章で設計を、第11章から第13章で実際の運用を詳しく見ていきます。

図2-2 Human Development Loop

Plan→Do→Check→Actで、人と組織の学習を週次で回す



観測する。しかし、監視しない

AIによって行動や本人の振り返りを継続的に扱う以上、一つ重要な境界があります。観測と監視は同じではありません。

本人とAIの対話には、本人が自分自身を振り返るための空間が必要です。すべての発言をそのまま上司が読むと分かっているならば、人は、自分が本当に考えていることではなく、「上司から正しく見える答え」を話すようになるでしょう。それでは学習にはなりません。

だから、本人の生の対話と、組織が学習するための情報を分けます。本人自身の内省空間は守る。一方で、前進や停滞、共通する阻害要因、チーム全体の傾向、構造上の問題など、組織が学ぶ必要のある情報は、必要な形へ構造化して返します。

もちろん、運用設計が不適切であれば、観測は監視へ変わります。よって、誰が何を見ることができるのか、何を評価に使うのか、何を本人だけの内省情報として扱うのか、その境界を明確にしなければなりません。

観測できることと、すべてを見てよいことは同じではありません。

この情報設計も、Human in the Loop の一部です。

人材成長循環(Human Development Loop)として 回す

ここまでを、企業経営でなじみのある計画・実行・確認・改善(PDCA)で整理すると、EvaliA の全体構造が見えます。

計画(Plan)では、本人のありたい姿を確認し、会社との接点を探し、チーム目標、Role、重要業績評価指標(KPI)を置きます。実行(Do)では、Role を具体的な期待行動へ落とし、本人が週次行動(Weekly Action)を選び、仕事の中で試します。

確認(Check)では、行動過程(Process)、心理状態(Mind)、成果(Result)を異なる窓から見ます。必要に応じて、チームの行動や共通する Gap も確認します。改善(Act)では、本人の次の行動を変えるのか、上司の関わり方を変えるのか、Role やチームの協力関係を変えるのか、それとも SCAD によって仕事の仕組みを変えるのかを判断します。そして、その結果を再び次の Plan へ戻します。

本書では、このように計画・実行・確認・改善(PDCA)を、人の成長、チーム学習、組織学習まで含む形へ具体化した循環を、**人材成長循環(Human Development Loop)**と呼びます。

PDCA そのものが新しいわけではありません。人材成長循環(Human Development Loop)とは、PDCA という基本的な運転原理を、**人と組織が学ぶ仕組みとして具体化したものです。**

図2-3

学習する組織を実装する全体構造

Learning Organization を目的に、EvaliA と SCAD で運転する構造整理

階層	名称	意味・役割
最上位目的	学習する組織 (Learning Organization)	個人・チーム・組織が、学び続けながら適応し、成長し続けることを目指す最上位目的
設計を検証する定規	五つのディシプリン (Five Disciplines)	学習する組織として設計が妥当かを点検するための理論的な定規
運転構造	人材成長循環 (Human Development Loop)	個人の学びと成長を、観測・対話・評価・再実行へと循環させる運転構造
実装エンジン	EvaliA	週次観測と対話を通じて、個人の状態・Gap・行動を可視化し、循環を実装するエンジン
組織改善サブループ	SCAD	個人の課題を、組織構造・役割・標準化・改善へ接続し、組織知へ変える補助ループ

Human in the Loop が貫く判断点

目標設定 ▶ 意味づけ ▶ Role設定 ▶ 行動選択 ▶ 評価 ▶ 介入 ▶ 組織知化

これらの判断点は AI に任せきるのではなく、人間参加型 (Human in the Loop) が一貫して担う。

AI は循環を支え、人間は意味づけと判断を担う。

Human in the Loop は、循環全体を貫く

では、この人材成長循環(Human Development Loop)を AI だけで回せばよいのでしょうか。私は、そうは考えていません。

AI が得意なのは、問いかけること、過去を記憶すること、変化を観測すること、複数の情報を比較すること、共通パターンを見つけること、別の視点を返すことです。

一方、人間に残すべきものがあります。本人が何をを目指すのか。何を大切にするのか。上司が誰にどの Role を任せるのか。本人が次に何を試すのか。チームとして何を優先するのか。成果をどう意味づけるのか。会社として、どの仕組みを変えるのか。そこには価値判断があり、責任があります。

よって、人工知能(AI)だけで完結させません。**人工知能(AI)が問い、観測し、比較し、整理する。人間が意味づけし、選び、判断し、責任を引き受ける。**これが、Human in the Loop です。

Human in the Loop は、AIが出した答えを最後に人間が承認するだけの仕組みではありません。本人の未来を決めるところ、Roleを決めるところ、AIの解釈を確認するところ、成果を評価するところ、どこへ介入するのかを決めるところ、個人の経験を組織知へ変えるところ。学習循環の重要な判断点に、人間が何度も参加します。

Human in the Loop とは、人間が意味と責任を手放さないための運転原則なのです。

本書の概念を、ここで一度整理する

ここまで多くの言葉が出てきましたが、それぞれの上下関係は明確です。

最上位の目的は、**学習する組織(Learning Organization)**です。その組織が、本当に学習する方向へ進んでいるのかを見る定規が、**五つのディシプリン(Five Disciplines)**です。日常業務の中で学習を回す具体的な循環が、**人材成長循環(Human Development Loop)**であり、その基本となる運転原理が**計画・実行・確認・改善(PDCA)**です。

その循環の判断点を貫き、人間が意味づけ、選択、判断、責任を担う運転原則が、**Human in the Loop** です。そして、この循環を日常的に回す実装エンジンが **EvaliA** です。さらに、循環の途中で組織構造上の問題が見つかったときに、仕事の標準、手順、制度、システムなどを改善するサブループが **SCAD** です。

この関係を混同してはいけません。EvaliA そのものが目的なのではありません。AIを導入することも目的ではありません。目的は、人と組織が学び続ける能力を持つことです。

二つのベクトルは、「一致させる」のではなく「調整し続ける」

二十年前の私は、「個人目標」と「組織目標」の方向が近づけば、組織力は高くなると考えていました。今も、その基本的な考え方は変わっていません。

ただし、現在は一つ重要なことを加えます。一致は固定された状態ではありません。本人が成長すれば、目指すものも変わります。会社の戦略も変わります。チームが変われば、本人に期待される Role も変わります。市場環境が変われば、昨日まで正しかった目標が、今日も正しいとは限りません。

重要なのは、一度完全に一致させることではありません。**ずれたことに気づき、ずれから学び、何度でも調整し直すこと**です。

AI の役割も、個人と会社の方向を無理に同じ方向へ向けることではありません。本人が何を目標しているのかを問い続ける。現在地を返す。会社とのずれを見えるようにする。行動の結果を返す。チームの状態を見る。仕組みの問題を見つける。そして、また本人と会社が選び直す。

AI は、一致を強制する装置ではなく、調整を止めないための装置です。

二十年前の図を、2026 年版へ

二十年前の図では、「個人目標」と「組織目標」という二つのベクトルが組織力へ向かっていました。現在は、その二つの方向の間に、学習の循環が入ります。

本人がやりたい姿を考え、会社との接点を探し、チームの中で Role を持つ。週次行動(Weekly Action)を選び、実行し、AI との対話で振り返り、次の行動を選ぶ。その途中で、本人だけでは変えられない問題が見つかればチームへ上げ、チームだけでは変えられないなら、SCAD によって組織の仕組みを変えます。

さらに、行動過程(Process)だけではなく、節目には心理状態(Mind)を振り返り、別の経路から成果(Result)を確認します。それらの間に Gap があれば、本人を直すと決めつけるのではなく、本人、上司、チーム、組織のどこを変えるべきかを考え、その結果を次の Plan へ戻します。

これが、人材成長循環(Human Development Loop)です。

二十年前の静的な「二つのベクトル」は、AI によって、**個人が学び、チームが学び、組織が学び、変わった組織の中でまた個人が学ぶ**という動的な循環へ変わりました。

では、「現在地」を何で見るのか

ここまでで、学習する組織(Learning Organization)を日常業務の中で運転する全体構造は見えてきました。しかし、一つ大きな問題が残ります。

本人へ、「あなたは、何になりたいのですか」と聞く。そして次に、「では、今のあなたは、そこに対してどこにいるのですか」と聞く。この二つ目の問いに、何を基準に答えればよいのでしょうか。

成長とは、知識が増えることでしょうか。技術が高くなることでしょうか。仕事が速くなることでしょうか。それとも、顧客への向き合い方、自律性、責任、周囲への貢献、他者を通じて成果を出す力まで含むのでしょうか。

私はこの問題を、「心」と「技」という二つのベクトルで考えてきました。ただし、二十年前と今では、そのモデルの意味づけは少し変わっています。以前は、人を一定の理想像へ育てるための地図として使っていました。現在は、本人が目指す未来に対して、「今、自分はどこにいるのか」を考えるための座標として読み直しています。

そして、その現在地を見るときには、もう一つ重要な問題が出てきます。自分が見ている現在地は、本当に現在地そのものなのか。それとも、自分のメンタルモデル(Mental Models)を通して見た現在地なのか。

次章では、人の成長を「心(Will)」と「技(Skill)」という二つの方向から捉えてきた人財マトリックスを、自己マスタリー(Personal Mastery)の「現在地」として読み直していきます。

第二部 人と組織の学習をどう設計するのか

第3章 人はどう成長するのか

——「心×技」の人財マトリックスを、自己マスタリー(Personal Mastery)から読み直す

第2章では、学習する組織(Learning Organization)を日常業務の中で動かす全体構造を見てきました。本人が何になりたいのかを考え、会社との接点を見つけ、チームの中で役割(Role)を持ち、小さな行動を起こし、その結果を振り返る。本人だけでは解決できない問題はチームや組織へ戻し、変わった組織の中でまた本人が学ぶ。この循環を繰り返すことが、人材成長循環(Human Development Loop)の基本です。

しかし、ここには一つ大きな問いが残っています。本人へ「あなたは何になりたいのですか」と聞くことはできます。では次に、「その未来に対して、あなたは今どこにいるのですか」と聞いたとき、何を基準に現在地を考えればよいのでしょうか。

知識が増えれば成長なのでしょう。仕事が速くなればよいのでしょうか。高い専門能力を持っていれば、その人は十分に成長した人材なのでしょう。それとも、顧客への向き合い方、自分で考える力、周囲への貢献、他者を通じて成果を出す力まで含めて見る必要があるのでしょうか。

私は、この問題を以前から「心」と「技」という二つの方向から考えてきました。それが、「人財マトリックス」です（できる若者は3年で辞める！50ページ参照）。

人の成長は、一つの物差しでは測れない

人の能力を考えると、最も分かりやすいのは「技」です。会計を知っているか、英語が使えるか、システムを操作できるか、営業経験があるか、ある仕事を一人で完結できるか。こうした知識、技術、経験は比較的観測しやすく、本人自身も不足に気づきやすいものです。

しかし、仕事で価値を生み出すためには、技だけでは足りません。どれほど専門能力が高くても、顧客の立場で考えなければ、知識が自分の中だけで完結してしまうことがあります。自分の担当業務を完璧に処理していても、チーム全体の成果には関心を持たない人もいます。高い専門性を持っていても、自分とは異なる意見を受け入れられなければ、周囲との協働を難しくすることもあります。

一方で、顧客の役に立ちたい、自分から動きたい、周囲へ貢献したいという意志は強くても、まだ十分な知識や経験を持っていない人もいます。この二人を、一つの「能力」という物差しだけで見れば、大切な違いが消えてしまいます。

そこで私は、人の成長を**心(Will)**と**技(Skill)**という二つのベクトルで考えるようになりました。拙著でも、人財マトリックスを「心」と「技」の二軸から人材育成の方向を考えるモデルとして示しています。

ここでいう「心」は、善人か悪人かという道徳的な評価ではありません。本書では、仕事への向き合い方や行動の方向を表す言葉として使います。顧客への関心、自律性、責任、周囲への貢献、他者の立場から考えようとする姿勢などです。

一方の「技」は、知識、技術、経験、判断、実行能力など、仕事を実際に遂行し、成果へ変えるための能力です。

ここで一つ、用語を明確にしておきます。本章でいう**心(Will)**は、**後に扱う心理状態(Mind)**とは別の概念です。心(Will)は、仕事への向き合い方や行動の方向を見るための人財マトリックス上の概念です。

一方、心理状態(Mind)は、本人が現在の仕事をどのように受け止めているのかという、満足、不満、意味、負担、成長実感などを見るための観測窓です。両者を混同してはいけません。

二十年前の人財マトリックス

私は、心と技という二つのベクトルから、人材の状態を四つに整理しました。「新人」「ネクスト・リーダー」「職人」「エース」です。

心と技の両方がまだ発展途上にある状態を「新人」としました。ここでいう新人は、年齢や入社年次ではありません。経験豊富な人であっても、新しい仕事や新しい Role へ移れば、その領域では再び新人になります。

心の方向は強いが、技がまだ発展途上にある状態が「ネクスト・リーダー」です。顧客や周囲への関心があり、自分から動こうとする。しかし、その意志を安定した成果へ変えるための知識や経験は、まだ十分ではない状態です。

反対に、技は高いものの、仕事への関心が自分の専門領域に集中し、顧客、チーム、組織への視野が相対的に弱くなっている状態を、私は「職人」と呼びました。そして、心と技の両方が高い状態を「エース」としました。

以前の著書では、この「新人」から「エース」へ進むことを、一つのキャリア・アップの方向として描いていました。これは、当時の私の人材育成思想をそのまま表したモデルです。

図 3-1

二十年前の人財マトリックス

— 「心」と「技」の二つのベクトルで見た人材の成長 —



二十年前の「Goal」を、今は少し違って見る

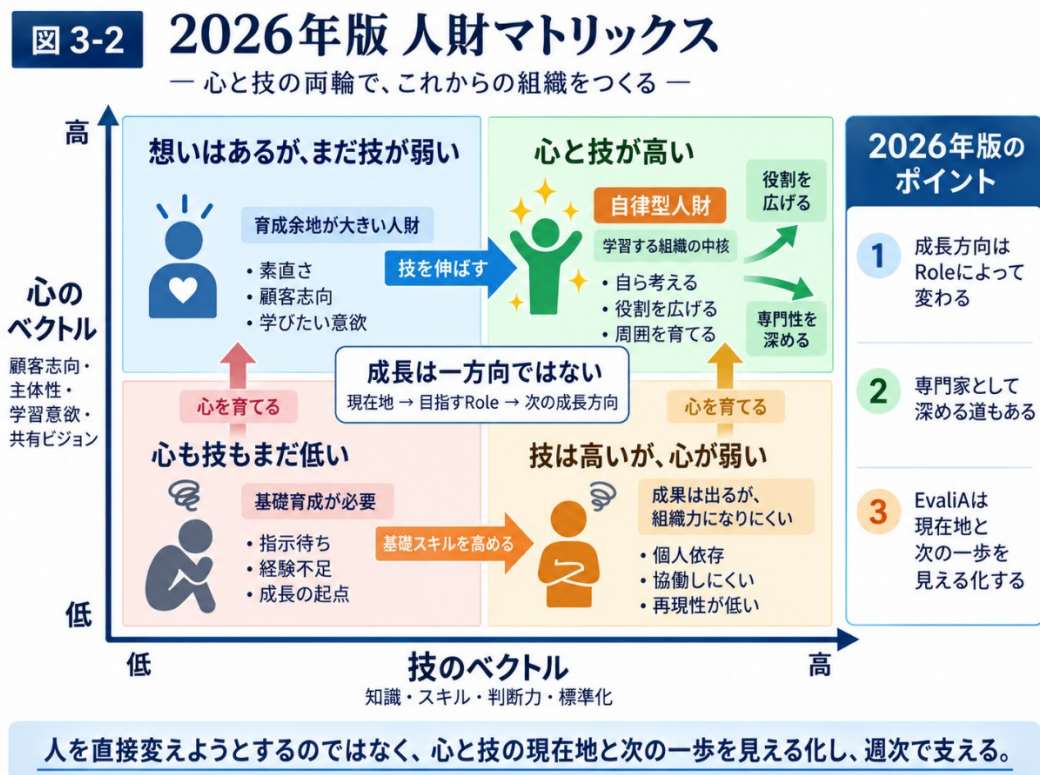
現在の私は、この人財マトリックスを捨てたわけではありません。しかし、使い方は変わりました。

二十年前の図では、「新人」から始まり、最終的に「エース」へ進むことを一つの成長方向として強く描いていました。それは、当時の私が、社員を一定の理想像へ育てることを重視していたからです。

現在は、そこを少し違って考えています。人には異なる役割(Role)があります。高度な専門職として価値を生み出す人もいます。すべての人が管理職になる必要はありません。また、同じ人であっても、目指す仕事や現在の Role が変われば、必要な心と技の内容も変わります。

したがって、現在の人財マトリックスで重視するのは、「あなたはどの象限の人間か」を決めることではありません。**現在の Role と、本人が目指す未来に対して、心と技のどちらに次の成長課題があるのか**を見ることです。

つまり、人財マトリックスを、人材を序列化する図から、成長の現在地を考える座標へ読み替えるのです。



人財マトリックスは、性格診断でも心理尺度でもない

ここで、人財マトリックスの位置づけを明確にしておく必要があります。

人財マトリックスは、人の性格を四種類に分類する診断ではありません。また、心理学的に確立された性格尺度をそのまま使っているわけでもありません。私が実務の中で、人材育成を考えるためにつくった二軸モデルです。

よって、人財マトリックスを見て、「この人は職人型だから、こういう性格だ」「この人はネクスト・リーダーだから、将来必ず管理職になる」と決めるべきではありません。

同じ人でも、仕事や Role が変われば現在地は変わります。新しい技術が登場すれば、それまで高い技を持っていた人でも再び学習が必要になります。優秀な専門職が初めて管理職になれば、専門技術では高い位置にいても、人を育てる技については学び始めたばかりかもしれません。

したがって、人財マトリックスの目的は「この人は何者か」を決めることではありません。「今どこにいて、次に何を学ぶ必要があるのか」を考えるための実務モデルです。

今後、十分なデータが蓄積されれば、信頼性や妥当性を含めて、より精緻な尺度として研究できる可能性はあります。しかし、それは今後の実証課題です。現段階での価値は、それまで定性的に語られることの多かった人の成長を二つの方向から構造化し、本人と上司が対話できるようにすることにあります。

技の不足は、比較の見えやすい

心と技には、一つ大きな違いがあります。自分自身で不足に気づきやすいかどうかです。

以前、私は、知識、技術、経験といった「技」の不足は、自分自身でも比較的認識しやすいと書いていました。英語を話せない人は、自分が英語を話せないことを知っています。難しい会計処理に直面すれば、自分に知識が足りないことが分かります。知らないシステムを操作できなければ、自分がその技術を持っていないことを認識できます。

不足が見えれば、本を読む、研修を受ける、資格を取る、経験者から学ぶ、実際の仕事を経験するといった学習行動へつなげやすくなります。

もちろん、技であれば必ず正確に自己認識できるという意味ではありません。専門家ほど自分を過大評価する場合がありますし、コミュニケーション能力のように、技術でありながら本人には不足が見えにくい領域もあります。

それでも、知識や専門技術には外部基準が存在する場合が多いため、心の領域と比べれば、「何が足りないのか」を認識しやすい傾向があります。

「心」の不足は、自分には見えにくい

難しいのは、心(Will)の方です。

例えば、「あなたは顧客志向ですか」と聞かれて、「いいえ。私は顧客のことを考えていません」と答える人は多くないでしょう。「あなたは責任感がありますか」と聞かれて、「私は責任感がありません」と答える人も少ないはずです。管理職へ「部下の話を聞いていますか」と聞けば、多くの人は「聞いている」と答えるでしょう。

しかし、本人がそう思っていることと、実際の行動が一致しているとは限りません。部下の話を聞いているつもりでも、途中で結論を言ってしまうかもしれません。顧客志向のつもりでも、実際には自社の商品説明に多くの時間を使っているかもしれません。主体的に動いているつもりでも、重要な判断になると必ず上司へ正解を求めていることもあります。

ここでは、単に「知らない」という問題とは違うことが起きています。自分が何を見えていないのか、そのこと自体が見えていないのです。

私はこれを、いわば「無知の無知」の状態と考えています。人が学び始めるためには、「自分にはここが見えていなかった」「自分の見方には前提があった」と気づく必要があります。つまり、無知の無知から「無知の知」へ移る必要があります。

ここで、メンタルモデル(Mental Models)が重要になります。

人財マトリックスとメンタルモデル(Mental Models)は、同じものではない

ここは明確に区別しておく必要があります。人財マトリックスの「心」を、そのままメンタルモデル(Mental Models)と考えるわけではありません。

心(Will)は、仕事への向き合い方や行動の方向を見ます。一方、メンタルモデル(Mental Models)は、その行動を生み出しているものの見方や前提を見ます。

例えば、ある管理職が部下へ仕事を任せないとします。人財マトリックスから見れば、「人を育てる」「他者を信頼して任せる」といった心の方向に課題がある可能性を考えることができます。

しかし、なぜ任せないのでしょうか。「部下へ任せれば失敗する」「上司は答えを持っていなければならない」「自分でやった方が顧客に迷惑をかけない」といった前提があるのかもしれませんが。この前提の方が、メンタルモデルです。

したがって、本書では、人財マトリックスは成長状態を見る座標であり、メンタルモデルは、その状態を生み出している認知や現在地の見え方を問い直すものとして分けます。

人財マトリックスは、自己マスタリー(Personal Mastery)そのものでもない

もう一つ、理論上の位置づけを明確にしておく必要があります。本書では、人財マトリックスを自己マスタリー(Personal Mastery)と接続して説明しています。しかし、人財マトリックスがセンゲの自己マスタリーそのものを測定する尺度だという意味ではありません。

自己マスタリー(Personal Mastery)は、本人が何を実現したいのかという方向を持ち、現在の現実を見つめ、その差から学び続ける、より広い概念です。

人財マトリックスが担うのは、その中の「現在地」を考える部分です。本人が目指す未来に対して、現在の心の方向はどうか。技の方向はどうか。どちらに次の成長課題があるのか。その問いをつくるための実務上の座標として使います。

整理すると、自己マスタリー(Personal Mastery)は、ありたい姿と現在地の差をつくる。人財マトリックスは、その現在地を実務上考えやすくする。そしてメンタルモデル(Mental Models)は、その現在地を本人がどう見ているかを問い直すという関係になります。

コミュニケーション能力は、「技」に見えて「心」 にも左右される

心と技を二つに分けると、「コミュニケーション能力はどちらに入るのか」という問題が出てきます。これは、二つの軸が完全には独立していないことを理解するうえで、分かりやすい例です。

質問の仕方、傾聴、要約、説明、フィードバック(Feedback)などには技術があります。訓練によって身につけることができます。その意味では「技」です。

しかし、実際のコミュニケーション行動は、その人のメンタルモデル(Mental Models)から強い影響を受けます。

例えば、コーチングや質問技法を学んだ管理職がいたとしても、「部下は経験が浅いからのだから、最後は自分が答えを教えるべきだ」と強く信じていれば、質問しているように見えても、最終的には自分の答えへ誘導するでしょう。

傾聴法の研修を受けていても、「自分が正しい答えを持っている」という前提が強ければ、相手の話を本当の意味で聞くことは難しくなります。この場合、さらに傾聴法を追加するだけでは十分ではありません。「なぜ自分は、すぐに答えを教えたくなるのか」と、自分の前提に気づく必要があります。

リーダーシップについても同じです。心と技は完全に独立した二つの箱ではなく、互いに影響し合います。だから人財マトリックスは、人を「心の人」「技の人」と分けるためのものではありません。成長を二つの方向から観測し、次に何を学ぶかを考えるために使うものなのです。

技では「追加学習」、心では「学び直し」が必要になることがある

心と技では、成長の仕方にも違いがあります。

技の成長では、新しい知識、技術、経験を追加する学習が中心になりやすいものです。知らないことを知る。できなかったことをできるようにする。経験を増やす。こうした学習です。

一方、心の領域では、新しいものを追加するだけでは行動が変わらないことがあります。「自分でやるのが一番速い」「顧客は価格しか見ていない」「部下へ任せれば失敗する」「上司が答えを教えるのが親切だ」といった前提を持っている場合、新しい知識を足しただけでは、以前と同じ行動を取り続けるかもしれません。

その場合、それまで正しいと思っていた見方そのものを問い直す必要があります。つまり、心の成長では、追加学習だけでなく、学習解除(Unlearning)やものの見方の組み替えが必要になる場合があります。

もちろん、これは絶対的な二分ではありません。高度な技術領域でも、過去の成功方法を捨てなければならないことがありますし、心の領域でも新しい知識を学ぶ必要があります。それでも、「知らないから教えればよい問題」と「本人が正しいと思っている前提そのものを問い直す必要がある問題」を分けて考えることには、大きな意味があります。

経理派遣事業で、なぜ「心を先に見る」ことが機能したのか

プロローグでは、経理人材の派遣事業が、私の人材育成への確信を強めたことを書きました。ここでは、その経験を人財マトリックスの観点から読み直してみます。

当時の派遣市場では、経験者や即戦力が求められることが一般的でした。一方、私たちは多くの未経験者にも門戸を開いていました。前掲書にも、「知識・技術・経験」だけでなく、「顧客に対する思い」を持てる人へチャンスを与えたいと考えていたことを書いています。

未経験者である以上、最初から技が高いわけではありません。しかし、技が低いことと、将来価値が低いことは同じではありません。顧客の立場で考えようとする。分からないことを素直に聞く。相手の役に立とうとする。自分の仕

事として責任を持とうとする。こうした仕事への向き合い方があれば、実務経験を積む中で技を高めていくことができます。

人財マトリックスで言えば、最初から高い技を持つことだけを採用条件にするのではなく、心の方向と、これから伸ばせる技を分けて見るという考え方です。この方法によって、未経験から仕事を覚え、顧客から評価される人が実際に現れました。その成功が、私に「心を先に見る」という考えへの強い確信を与えました。

しかし、現在はこの経験も、もう少し慎重に捉えています。「心があれば技はなくてもよい」という意味ではありません。仕事によっては、一定の技を持っていなければ任せてはいけないものがあります。高度な専門判断や、安全に重大な影響を与える仕事では、必要な知識や技術を満たすことが前提になります。

人財マトリックスが示しているのは、技が重要ではないということではありません。**技だけでは、人の成長を十分に説明できない**ということです。

「職人」は、悪い人材なのか

ここで、人財マトリックスの言葉について一つ誤解を避けておく必要があります。

本モデルでいう「職人」は、日常語でいう熟練職人を否定する言葉ではありません。また、高度な専門職が管理職より劣るという意味でもありません。むしろ、高い専門性は会社にとって極めて重要です。すべての専門家が管理職になる必要もありません。

問題は専門性が高いことではなく、その専門性が自分の中だけで完結してしまうことです。

例えば、非常に優秀な専門家が技術的には完全に正しい回答を出せるとします。しかし、顧客が本当に困っている問題が別にあるなら、その技術的な正解だけでは十分な価値にならないことがあります。また、自分だけができる仕事を増やし続け、後輩やチームへ知識を移さなければ、その専門性は会社にとって属人的なリスクにもなります。

したがって、「職人」から「エース」へ進むという従来の表現を、現在は単純に「専門家から管理職へ進む」とは考えていません。

高度な専門職であっても、顧客価値を考え、自分の知識をチームへ還元し、組織の成果へつなげることができれば、心と技の両方が強い状態と考えることができます。

つまり、現在の意味での「エース」は、管理職の別名ではありません。**自分の技を、自分以外の価値へ変換できる人**です。

「ネクスト・リーダー」は、心だけで完成しているわけではない

反対に、「ネクスト・リーダー」も完成した人材ではありません。顧客志向が強く、主体性があり、周囲への関心もある。しかし、技がまだ十分ではない。その人には高い成長可能性があります、現時点では任せられない仕事もあります。

良い意図があっても、知識が足りなければ正しい判断はできません。責任感が強くても、技術が不足していれば成果にはつながりません。

したがって、ネクスト・リーダーに必要なのは、さらに精神論を教えることではありません。その強い意志を成果へ変えるために、次にどの知識、技術、経験を身につける必要があるのかを明確にすることです。

ここでも、人財マトリックスは人を評価する表ではなく、現在地から次の学習方向を考える座標として機能します。

エースになれば、成長は終わるのか

では、心と技の両方が強い「エース」になれば、成長は終わるのでしょうか。もちろん、そうではありません。

自分ができるようになった後には、別の学習が始まります。自分の暗黙知を言語化する。他の人ができるようにする。成功した方法の中から再現できる部

分を見つける。自分が得た知識をチームへ共有する。必要なら標準や AI へ戻す。

ここまで進んで初めて、個人の能力が組織の能力へ変わります。

前掲書では、心と技の両方を持つエース領域をキャリア・アップの一つの方向として示していました。しかし現在は、その先まで考えています。

エースは個人としての完成ではなく、個人の学びを組織の学びへ変える出発点でもある。

自分の知識を共有し、他者へ渡し、必要に応じて標準や AI へ反映することができれば、一人の学びが会社全体の能力へ広がります。この問題は、後に扱うチーム学習(Team Learning)や名人循環(Mastery Loop)へつながっていきます。

人財マトリックスは、「点数表」ではない

AI や人事システムを使うと、人を数値化したくなる誘惑があります。心が 72 点、技が 84 点、だからこの人はこの象限にいる、と表示すれば、非常に客観的に見えます。

しかし私は、人財マトリックスをそのような単純な点数表として扱うべきではないと考えています。

心(Will)は、直接測ることのできる物理量ではありません。顧客志向、自律性、責任感、他者視点などは、行動から推測することはできます。しかし、「心そのもの」を直接測定することはできません。

したがって、現時点で求めるべきなのは、精密な数字に見せることではありません。定性的だった人の成長を二つの方向に構造化し、本人と上司が観測し、対話できるようにすることです。

今、自分はどちらの方向が強いのか。どちらに次の成長課題があるのか。本人と上司では現在地の見方が違ってないか。数か月前と比べて、どちらの方向へ動いているのか。そうした問いを生み出すことに、人財マトリックスの価値があります。

人財マトリックスを、自己マスタリー(Personal Mastery)の現在地として読み直す

ここまで見てくると、人財マトリックスの位置づけは、二十年前とはかなり違って見えてきます。

以前の私は、人財マトリックスを主に「人材育成の方向を示すモデル」として使っていました。現在は、それを、自己マスタリー(Personal Mastery)に必要な「現在地」を実務上考えるための座標として読み直しています。

本人には「何になりたいのか」という未来があります。しかし、未来だけでは学習は始まりません。現在どこにいるのか、何が足りないのか、次にどちらの方向へ進む必要があるのかを考える必要があります。

技の不足が明確なら、知識、技術、経験を増やすという課題が見えます。一方、心の側で停滞している場合には、単に研修を増やしても動かないことがあります。なぜ、その行動が起きているのか。本人は何を当然だと思っているのか。何を恐れているのか。どのような成功体験が、その見方を強くしているのか。ここで、メンタルモデル(Mental Models)を見る必要が出てきます。

整理すると、自己マスタリー(Personal Mastery)は、本人が目指す未来と現在地の差をつくります。人財マトリックスは、その現在地を心(Will)と技(Skill)という二つの方向から考えやすくします。そしてメンタルモデル(Mental Models)は、その現在地を本人自身がどのような前提で見ているのかを問い直します。これらを組み合わせることで、人の成長をより立体的に捉えることができます。

人を一つの物差しで決めない

この章で最も伝えたいのは、「心の方が技より重要だ」ということではありません。まして、「エースが偉く、職人が劣っている」ということでもありません。

重要なのは、人を一つの物差しだけで決めないことです。技だけを見れば、専門性の高い人が最も優秀に見えるかもしれません。心だけを見れば、意欲の高い人を過大評価するかもしれません。しかし、人の現在地は一つの情報だけでは分かりません。

まず心と技という二つの方向から見る。そして、本人自身の認識も絶対視しない。本人が見ている現在地と、周囲から見える行動との間に違いがあれば、その違い自体を学習材料にします。ここから、第二部全体を貫く一つの原則が始まります。

一つの数字、一つの行動、一つの評価、一つの視点だけで、人を決めない。複数の視点を持ち、その間にある違いから学ぶ。人財マトリックスは、その最初の入口です。

現在地が分かっても、「自分にはどう見えているか」が残る

人財マトリックスによって、心と技という二つの方向から現在地を考えることはできます。しかし、それでも問題は残ります。

本人自身が、その現在地を正しく見ているとは限らないからです。

周囲から見れば、顧客への関心が弱く見えている。しかし本人は、「自分は十分に顧客志向だ」と思っているかもしれません。上司から見れば、ほとんど部下へ任せていない。しかし本人は、「かなり任せている」と考えているかもしれません。

そうであれば、成長課題を示すだけでは学習は始まりません。なぜ、自分にはそう見えているのか。自分は、どのような前提から現在地を見ているのか。そこまで考える必要があります。

人財マトリックスで「現在地」を考えた次に必要になるのは、**その現在地を見る自分自身のレンズを問い直す**ことです。

次章では、自分には見えにくい自分自身の前提を、メンタルモデル(Mental Models)としてどのように学習対象へ変えるのかを考えていきます。

第4章 自分の現在地は、なぜ自分には見えないのか

——メンタルモデル(Mental Models)と自己マスタリー(Personal Mastery)

第3章では、人の成長を「心(Will)」と「技(Skill)」という二つの方向から見る人財マトリックスについて考えました。人財マトリックスは、人を固定的に分類するためのものではありません。本人が現在どこにいて、目指す未来や現在の役割(Role)に対して、次にどの方向へ成長する必要があるのかを考えるための実務上の座標です。

しかし、座標があれば現在地を正確に見ることができるとは限りません。なぜなら、本人自身が、自分の現在地を正しく見ているとは限らないからです。

周囲から見れば、顧客の話を十分に聞けていない。しかし本人は、「私は顧客志向だ」と考えている。部下から見れば、ほとんど仕事を任せてもらえていない。しかし上司は、「かなり任せている」と思っている。本人には、嘘をついているつもりはありません。本当に、そのように見えているのです。

技術や知識の不足であれば、「知らない」「できない」と本人自身が認識しやすい場合があります。しかし、自分がどのような前提でものごとを見ているのかは、自分自身には見えにくいものです。

この「ものの見え方」を扱うのが、メンタルモデル(Mental Models)です。

自己マスタリー(Personal Mastery)とメンタルモデル(Mental Models)は何が違うのか

自己マスタリー(Personal Mastery)とメンタルモデル(Mental Models)は、どちらも本人の内面に関係するため、混同しやすい概念です。本書では、実務上の役割を明確に分けて考えます。

自己マスタリー(Personal Mastery)は、「私はどこへ行きたいのか。そして、その未来に対して現在どこにいるのか」を扱います。

一方、メンタルモデル(Mental Models)は、「私は、その現在地や周囲の世界を、どのような前提を通して見ているのか」を扱います。

例えば、ある社員が「将来は海外拠点の経営を任されたい」と考えているとします。本人には財務知識も顧客対応経験もある。しかし、人を育て、自分がいなくてもチームとして成果を出せる状態をつくる経験は、まだ十分ではない。この場合、ありたい姿と現在地の間には差があります。その差を認識し、次に何を学ぶ必要があるのかを考えるのが自己マスタリーです。

ところが本人が、「自分は人を育てることも十分できている」と考えていたらどうでしょうか。あるいは、「部下のためには、自分が答えを教えてあげる方がよい」と当然のように考えていたらどうでしょうか。この場合、単に能力が不足しているだけではありません。現在地そのものの見え方に、本人の前提が入り込んでいます。

したがって、本書では、**自己マスタリー(Personal Mastery)は、ありたい姿と現在地の差をつくる。メンタルモデル(Mental Models)は、その現在地をどう見ているのかを問い直す。**このように整理します。

人は「現実そのもの」ではなく、「自分が解釈した現実」を見ている

人は、何の前提も持たずに世界を見ることはできません。過去の経験があります。成功体験も失敗体験もあります。以前の上司から教えられたこと、所属してきた会社の文化、これまで自分を守ってきた考え方もあります。私たちは、それらを通して目の前で起きている出来事を解釈しています。

例えば、若い頃にミスをすると厳しく叱責される職場で育った人は、管理職になったとき、部下へ重要な仕事を任せることに強い不安を感じるかもしれません。過去に部下へ任せて大きな失敗を経験した人なら、「重要な仕事は、最後には自分で確認しなければならない」と考えるでしょう。営業として「訪問件数を増やせば売上が上がる」という経験で成功してきた人なら、環境が変わった後も「売上が足りないなら、もっと訪問すればよい」と考えるかもしれません。

本人にとって、こうした考えは単なる思い込みではありません。多くの場合、自分自身の経験から導き出した合理的な結論です。だからこそ、メンタルモデルは自分自身には見えにくくなります。

「私は、一つの見方をしている」と認識していれば、別の見方を考えることができます。しかし、「これは現実そのものだ」と思っていれば、自分の見方を問い直す必要を感じません。

メンタルモデルを学習対象にする第一歩は、自分の考え方を否定することではありません。**自分にも、自分なりの見方があると気づくことです。**

成功体験ほど、強いメンタルモデル(Mental Models)になる

メンタルモデルは、失敗からだけ生まれるものではありません。むしろ、成功から生まれたものの方が、強く固定されることがあります。

私自身が、その例でした。私は、自分の考え方を疑い、努力を続け、行動を修正することで経営上の壁を越えてきました。その方法が、私自身には有効でした。そして、同じような教育によって大きく変わる社員も実際にいました。すると、「やはり、人は考え方を換えれば成長できる」という確信がさらに強くなります。成功が、自分の前提を強化したのです。

ここに、成功体験の難しさがあります。失敗した方法なら捨てやすい。しかし、成功した方法は手放しにくい。過去に正しかったからこそ、「今も正しいはずだ」と考えてしまいます。

これは経営にも、人材育成にも起こります。過去に細かく管理することで成果が出た上司は、部下の能力が上がった後も同じ管理を続けるかもしれません。自分が長時間働くことで成功した経営者は、次の世代にも同じ働き方を求めるかもしれません。

しかし、環境や役割(Role)が変われば、過去の成功方法が現在の成長を止めることがあります。したがって、人が学ぶためには、失敗を振り返るだけでは足りません。**自分を成功させた前提が、現在も有効なのかを問い直すことも必要なのです。**

問題は、「知らないこと」より「知らないことが見えないこと」にある

第3章では、「無知の無知」から「無知の知」へ移ることが重要だと述べました。

英語が話せないことは、自分でも比較的分かります。会計知識が不足していれば、難しい問題に直面したときに気づくことができます。

しかし、「私は本当に顧客の立場で考えているか」「私は本当に部下へ任せているか」「私は自分の責任として問題を考えているか」といった問いには、簡単な外部基準がありません。

しかも本人には、自分なりの説明があります。「顧客のためを思って、あえて厳しく言っている」「部下を守るために、自分が最後に確認している」「会社の仕組みが整えば、自分ももっと動ける」。どれも、一部は正しい可能性があります。

上司が「あなたの考え方は間違っています」と結論づけても、学習は起きにくいのです。必要なのは、正しい価値観を外から教えることではありません。本人自身が、「別の見方もあるかもしれない」と気づくことです。

自分の考え方を一段上から見る。自分が何を当然だと思っているのかに気づく。この働きは、メタ認知(Metacognition)とも深く関係しています。ここに、AIとの対話を使う意味があります。

AI は、「教師」ではなく「鏡」になる

AIによる人材育成というと、AIが社員へ正しい答えを教える姿を想像するかもしれませんが、本書では、それを中心には置きません。

AIが「あなたの問題はこれです」「あなたはこう考えるべきです」と決めてしまえば、従来の上司による価値観教育を、AIへ置き換えただけです。それでは、プロローグで述べた問題を別の形で繰り返すことになります。

重要なのは、答えを与えることではなく、本人には見えていなかった差を返すことです。

例えば本人が、「私は部下へ任せています」と話している。一方、過去数週間の振り返りを見ると、重要な判断については毎回、自分で最後の答えを出している。このときAIは、「任せているという自己認識と、実際に話してきた判断行動との間に差がある可能性はありませんか」と問い返すことができます。

あるいは本人が三週間続けて、「時間がなくてできませんでした」と話している。過去の振り返りを見ると、毎週、緊急業務を優先し、重要だが緊急ではない仕事を後回しにしている。そこでAIは、「時間そのものの不足だけでなく、優先順位のつけ方が関係している可能性はありませんか」と返すことができます。

AIが本人の心を読んでいるわけではありません。本人が語ったこと、与えられた情報、許された範囲で観測された行動や変化をもとに、別の見方を仮説として返しているだけです。

本書では、この働きを**内省支援フィードバック(Reflective Feedback)**と呼びます。

AIの価値は、「時間軸を持った鏡」になれることに ある

AIを鏡として使う価値は、良い質問を一回できることだけではありません。過去の自分と現在の自分を比較できることにあります。

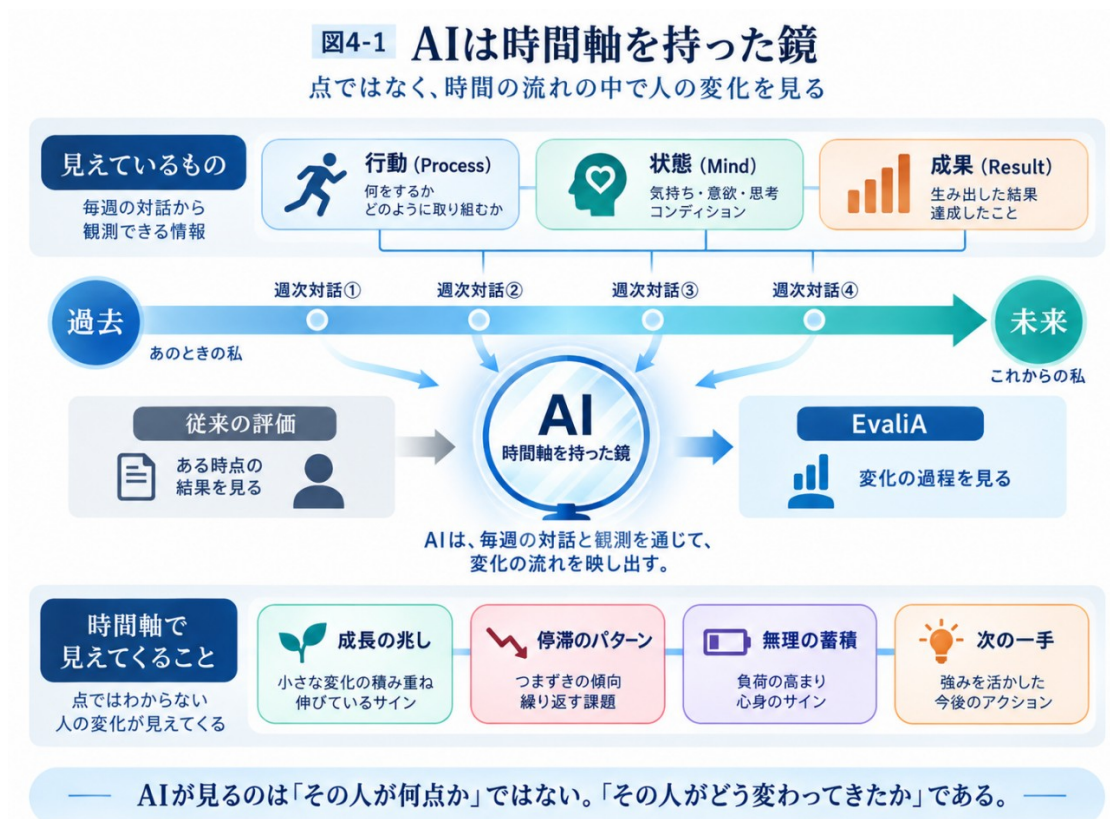
人の自己認識は、そのときの出来事や感情に強く影響されます。今週うまくいけば、「かなり成長した」と感じるかもしれません。一度大きく失敗すれば、「自分は何も成長していない」と感じることもあります。しかし、人の成長は一週間だけを切り取っても分かりません。

一か月前には、上司へ「どうすればよいですか」と答えを求めていた人が、現在は「私はこう考えますが、どうでしょうか」と自分の案を持って相談するようになっているかもしれません。三か月前には、顧客との面談で自社の商品説明から始めていた人が、現在は相手が何を実現したいのかを聞いてから話すようになっているかもしれません。

反対に、本人は「成長している」と感じていても、三か月間、同じ理由で同じ行動を先送りしていることもあります。

つまり、現在の本人を一点だけで見るのではなく、**時間の中でどちらへ動いているのかを見る**必要があります。

AIは、瞬間を映す鏡だけではなく、時間軸を持った鏡になることができます。これは、自己マスタリー(Personal Mastery)にも重要です。「今どこにいるか」だけではなく、「どちらへ動いているか」まで見るができるからです。



視点を変えると、問題の意味が変わる

メンタルモデル(Mental Models)を問い直すためには、時間軸だけでなく、視点を増やすことも重要です。

例えば、部下が予定どおり仕事を終えられなかったとします。上司から見れば、「もっと早く相談してほしかった」と感じるかもしれません。部下から見れば、「上司が忙しそうで相談できなかった」と考えているかもしれません。顧客から見れば、「誰に原因があるかより、早く問題を解決してほしい」と思うでしょう。チーム全体から見れば、「この仕事が一人に集中しすぎていること」が原因かもしれません。

さらに時間軸を広げれば、毎月同じ時期に同じ問題が起きている可能性もあります。

視点を変えると、「部下の能力不足」に見えていた問題が「上司との関係」の問題へ変わり、さらに「仕事の設計」の問題へ変わることがあります。

ここで、メンタルモデル(Mental Models)とシステム思考(Systems Thinking)がつながります。本人から見たらどうか。上司から見たらどうか。顧客から見たらどうか。他部署から見たらどうか。短期と長期ではどう違うか。全員が同じ行動をしたら何が起こるか。この問題は過去にも繰り返されていないか。

こうした問いを日常の仕事の中で繰り返すことによって、システム思考を単なる研修科目ではなく、仕事の見方そのものへ変えることができます。

メンタルモデル(Mental Models)を持っているのは、部下だけではない

人材育成というと、部下の考え方をどう変えるかに意識が向きやすくなります。しかし、メンタルモデル(Mental Models)を持っているのは部下だけではありません。上司にもあります。

「この社員にはまだ任せられない」「若い社員には判断できない」「営業は数字がすべてだ」「この仕事は昔からこの方法でやってきた」「この社員は主体性がない」。こうした前提は、上司自身の Role 設定やフィードバック(Feedback)、評価に影響します。

例えば本人は、「もっと仕事を任せてほしい」と考えている。一方、上司は「まだ任せられない」と考えている。このとき、どちらか一方を正解にしても問題は解決しません。

何を任せたのか。どこまで本人が判断したのか。どのような結果になったのか。上司はどこまで介入したのか。本人が失敗できる範囲はどこまでだったのか。具体的な事実を見る必要があります。

第3章では、「心」と「技」という複数の方向から現在地を見ました。本章ではさらに、本人、上司、顧客、チーム、時間軸という複数の視点から同じ現象を見ます。

ここでも貫く原則は同じです。**一つの情報、一つの評価、一つの視点だけで、人を決めない。**

「正しい価値観」を教えることから離れる

これは、私自身の人材育成にとって大きな転換です。以前の私は、社員には「正しい考え方」を持ってもらう必要があると考えていました。顧客志向であるべきだ。主体的であるべきだ。自分の責任として問題を考えるべきだ。今でも、これらは仕事をするうえで重要な方向だと思っています。

しかし、「だから、あなたもそう考えるべきです」と教えることと、「あなたが今取っている行動を別の立場から見ると、どう見えますか」と問いかけることは違います。

前者では、教える側が答えを持っています。後者では、本人が答えを探します。

学習する組織(Learning Organization)をつくるのであれば、中心に置くべきなのは後者です。人材育成とは、上司の価値観を部下へコピーすることではありません。人工知能(AI)が会社の価値観を社員へ植え付けることでもありません。

本人が、自分のありたい姿、実際の行動、周囲との関係を材料にして、自分自身の見方を更新できる状態をつくる。ここで、人材育成は「人を直接変えること」から、「人が自ら学ぶ条件をつくること」へ変わります。

学習解除(Unlearning)は、過去の自分を否定することではない

自分の前提を問い直すというと、「これまでの考え方は間違っていた」と認めなければならないように聞こえるかもしれませんが、しかし、過去の前提が間違っていたとは限りません。

例えば、起業直後には「自分でやった方が速い」という考え方は合理的です。社員が数人しかいなければ、経営者自身が判断し、実行した方が会社は速く動きます。

しかし、社員が100人になっても経営者がすべてを自分で判断していたら、今度は経営者自身が会社の成長を止めるボトルネックになります。

新人時代には、「分からなければ上司へ聞く」が適切です。しかし管理職になれば、すべての判断を上司へ求め続けることは適切ではなくなります。専門職では「自分が完璧にできる」ことに大きな価値があります。しかし管理職になれば、「他の人ができるようにする」ことの価値が高まります。

つまり、役割(Role)や状況(Context)が変われば、以前は有効だった前提が、現在には合わなくなることがあります。

学習解除(Unlearning)とは、過去の自分を否定することではありません。過去には有効だった前提を、現在の Role や Context に合わせて選び直すことです。

メンタルモデル(Mental Models)を「変える」のではなく、「選べる」ようにする

ここまで考えると、メンタルモデル(Mental Models)を学ぶ目的も見えてきます。一つの「正しいメンタルモデル」へ本人を変えることはありません。状況に応じて、複数の見方から選べるようになることです。

例えば、「部下へ任せれば失敗する」という考えにも、適切な場面があります。重大事故につながる仕事を、経験のない社員へいきなり全面的に任せるべきではありません。

一方、育成を目的とする仕事では、一定の失敗可能性を受け入れなければ、本人に判断経験は蓄積されません。つまり、「任せることが正しい」「任せないことが正しい」という二者択一ではありません。状況(Context)によって適切な選択は変わります。

重要なのは、自分は今どの前提を使っているのか。その前提は現在の状況に合っているのか。別の前提を選ぶことはできないのか、と考えられることです。

メンタルモデルをなくすのではなく、自分のメンタルモデルに支配されず、必要に応じて選び直せるようになる。私は、そこに成長の大きな意味があると考えています。

安心して内省できる情報設計が必要になる

ただし、自分の前提を率直に振り返るには条件があります。

例えば本人が、「この AI との対話を上司がすべて読み、そのまま人事評価に使う」と考えていたらどうなるでしょうか。「今の上司には不満があります」「会社の方向に納得できていません」「現在の仕事に意味を感じられません」といった本音は、話しにくくなるでしょう。代わりに、評価されやすい答えを選ぶようになります。

それでは、メンタルモデル(Mental Models)を振り返るための対話にはなりません。

そこで必要なのが、安心して内省できる情報設計です。本人との生の対話は、本人自身が考えるための空間として扱う。

一方、チームや組織が学ぶ必要がある場合には、個々の生の発言をそのまま渡すのではなく、共通する阻害要因や行動傾向、構造的な問題のシグナルなど、組織学習に必要な形へ整理して返す。

本人が率直に自分を振り返れることと、組織が必要な問題から学べること。この二つを両立させる必要があります。これは単なる情報管理の問題ではありません。**学習が本当に起こるかどうかを左右する設計上の問題**なのです。

AI の鏡にも、歪みはあり得る

AI を鏡として使うのであれば、鏡そのものも常に正しいわけではないことを忘れてはいけません。AI が返す問いやフィードバック(Feedback)は、与えられた情報、設定された目的、会社側から与えられた Role や方針、AI 自体の特性によって変わります。

「AIがそう言ったから正しい」としてはいけません。AIができるのは、本人の発言や行動、過去からの変化をもとに、一つの可能性を示すことです。本人の内面を直接知ることはできません。

したがって、AIが返すものは、判定ではなく仮説として扱います。「あなたの原因はこれです」と断定するのではなく、「こういう可能性はありませんか」と問い返す。本人が考え、必要なら実際の行動や周囲の事実とも照らし合わせる。最終的な意味づけは、人間が行います。

自己マスタリー(Personal Mastery)は、メンタルモデル(Mental Models)によって深くなる

ここまで見てくると、自己マスタリー(Personal Mastery)とメンタルモデル(Mental Models)の関係が、より明確になります。

自己マスタリーは、「私は何になりたいのか」という方向をつくります。そして、人財マトリックスなどを使い、「今、自分はどこにいるのか」という現在地を考えます。

しかし、その現在地の認識にも、自分自身の前提が入っています。

ありたい姿を持ち、現在地を見た後に、自分がどの前提から現在地を見ているのかを問い直す必要があります。別の視点からもう一度現在地を見て、ありたい姿との差を捉え直し、次の行動を選びます。

つまり、自己マスタリー(Personal Mastery)が成長の方向と現在地の差をつくり、メンタルモデル(Mental Models)が**その現在地を見る解像度を高める**という関係です。

この二つが組み合わさることで、単なる能力開発から、「自分自身の学び方を学ぶ」段階へ進むことができます。

自分を知っただけでは、学習する組織(Learning Organization)にはならない

しかし、ここまでの話には、まだ足りないものがあります。本人が何になりたいのかが分かった。現在地も見えてきた。自分がどのような前提でものを見ているのかにも、少しずつ気づけるようになった。それでも、ここまでで成立しているのは個人学習です。

会社は一人では成り立ちません。本人が目指す未来と、会社が目指す未来をどこでつなぐのか。自分が成長したい方向と、チームが現在必要としているものが異なったらどうするのか。自分が得意な行動だけ続けるのではなく、チームとして不足している機能を誰が担うのか。

ここから、個人の学習をチーム学習(Team Learning)へ進める必要があります。

その接点になるのが、共有ビジョン(Shared Vision)、役割(Role)、そしてCAGE 適応モデル(CAGE Adaptation Model)です。

人財マトリックスが、「私は今どこにいるのか」を考える座標だとすれば、CAGE 適応モデルは、**「このチーム、この状況の中で、私はどのような Role と行動を選ぶ必要があるのか」**を考えるための実務モデルです。

自分の現在地を知るだけでは、組織にはなりません。次に必要なのは、自分の未来と、チームや会社の未来を接続することです。

次章では、共有ビジョン(Shared Vision)を単なる理念の共有で終わらせず、チーム目標、CAGE 適応モデル、役割(Role)へ落とし、個人の学びをチーム学習(Team Learning)へ変える方法を考えていきます。

第5章 個人の未来を、チームと会社の未来へどうつなぐか

——共有ビジョン(Shared Vision)・役割(Role)・CAGE 適応モデル(CAGE Adaptation Model)

第4章では、自己マスタリー(Personal Mastery)とメンタルモデル(Mental Models)について考えました。本人が何になりたいのかを考え、自分の現在地を見つめ、さらに「自分はその現在地を、どのような前提を通して見ているのか」を問い直す。ここまで進めば、個人としての学習はかなり深くなります。

しかし、それだけでは学習する組織(Learning Organization)にはなりません。一人ひとりが自分の未来だけを追求し、それぞれが別の方向へ進めば、組織として一つの力にはならないからです。反対に、会社が自分たちの未来だけを社員へ押しつけても、人は自律的には動きません。

本人はどこへ行きたいのか。会社はどこへ行こうとしているのか。その中でチームは何を実現しなければならないのか。そして、一人ひとりは何を担うのか。この四つをつなげる必要があります。

その接点になるのが、共有ビジョン(Shared Vision)、役割(Role)、そしてCAGE 適応モデル(CAGE Adaptation Model)です。本章では、個人の学習をどのようにチームの学習へ変え、さらに会社が目指す未来へつないでいくのかを考えていきます。

共有ビジョン(Shared Vision)は、会社のビジョンを覚えることではない

「共有ビジョン」という言葉を聞くと、会社の経営理念やビジョンを社員全員が理解し、同じ方向を向くことだと考えるかもしれません。もちろん、会社がどこへ向かっているのかを社員が理解することは必要です。

しかし、会社が立派なビジョンをつくり、それを社員へ説明し、社員が内容を暗記したからといって、それだけで共有ビジョン(Shared Vision)が生まれるわけではありません。そこには、本人自身の未来が入っていないからです。

第2章で述べたように、私は二十年ほど前から、組織を「個人目標」と「組織目標」という二つのベクトルで考えてきました。本人は何をしたいのか。会社は何を実現したいのか。この二つの方向に接点があるほど、会社の仕事と本人の成長はつながりやすくなります。

本書では、共有ビジョンを実務へ落とし第一歩として、個人が目指す未来と会社が目指す未来との接点を、本人自身が見つけることを重視します。

例えば、ある社員が「将来は海外拠点を経営できる人材になりたい」と考えていて、会社も海外事業の拡大を目指しているのであれば、二つの方向には明確な接点があります。本人にとって海外案件へ関わることは、単に会社から命じられた仕事ではなく、自分自身の未来へ近づく経験になります。一方、会社にとっても、その人の成長が組織の未来につながります。

ここで初めて、会社の未来と個人の未来が同時に動き始めます。

共有できない部分まで、共有したことにしない

ただし、個人と会社の未来は、必ずしも完全には一致しません。

例えば本人は、「専門家として一つの分野を深く追究したい」と考えている。しかし会社は、「今後は管理職として、多くの人を率いてほしい」と考えている。ここには明らかなずれがあります。

このとき、「会社のビジョンに共感してください」と伝えても、ずれそのものは消えません。むしろ、違いがあるのに、それをないことにしてしまう方が危険です。

現在は接点が小さいだけかもしれません。Roleを変えれば接点が広がるかもしれません。対話を続ける中で、本人自身の目標が変わることもあれば、会社の側が「この人には別の活躍の仕方がある」と考え直すこともあるでしょう。場合によっては、将来的に本人と会社が別の方向へ進むこともあります。

共有ビジョン(Shared Vision)とは、全員がまったく同じ夢を持つことではありません。どこが重なり、どこが重なっていないのかを理解したうえで、共に目指せる方向を見つけることです。

共有できないものまで「共有している」と扱えば、会社は社員の本当の考えを見失います。だから私は、二つの方向のずれそのものも、学習材料として扱う必要があると考えています。

個人と会社の接点を、「チーム目標」へ落とす

個人と会社の未来に接点が見つかって、それだけでは日常の仕事にはなりません。会社のビジョンは、一人ひとりの日々の行動を決めるには大きすぎるからです。

「海外事業を拡大する」という会社の方向があったとしても、それだけでは一人の社員が明日何をすればよいのかは分かりません。そこで、会社のビジョンや戦略を、部門やチームが現在実現すべき目標まで具体化します。

例えば、「東南アジアで新規顧客を開拓する」「海外拠点の管理品質を標準化する」「現地マネージャーを育成する」といった目標です。ここで初めて、会社の未来が「このチームとして、今何を実現するのか」というレベルまで近づいてきます。

しかし、チーム目標が決まっても、まだ十分ではありません。次に考えるべきなのは、その目標を実現するために、**このチームにはどのような機能が必要なのか**ということです。

この問いを考えるために、本書で使うのが CAGE 適応モデル(CAGE Adaptation Model)です。

CAGE 適応モデル(CAGE Adaptation Model)——

チームに必要な四つの機能

CAGE 適応モデルでは、仕事の中で必要となる行動機能を、**成果(Challenger)、探究(Adventurer)、安全(Guardian)、関係(Empathizer)**という四つの方向から考えます。

成果(Challenger)は、目標を達成し、結果に責任を持ち、必要な決断を行い、困難があっても前へ進める機能です。

探究(Adventurer)は、新しい可能性を探し、問いを立て、未知の方法を試し、既存の前提を疑う機能です。

安全(Guardian)は、計画し、確認し、リスクを管理し、品質や継続性を守る機能です。

関係(Empathizer)は、相手を理解し、支援し、信頼をつくり、人と人をつないで協働できる状態をつくる機能です。

CAGE では、この四つを二つの軸で整理します。横軸は「自律と共生」、縦軸は「変化と安定」です。

変化・新しさと共生が重なる領域に成果(Challenger)、変化・新しさと自律が重なる領域に探究(Adventurer)、安定・安全と自律が重なる領域に安全(Guardian)、安定・安全と共生が重なる領域に関係(Empathizer)を置きます。

ここで、このモデルの位置づけを明確にしておきます。CAGE 適応モデルは、確立された性格診断や心理尺度として提示するものではありません。本書では、私たちが実務と研究の中で整理しているモデルとして、**Role と Context に応じて、チームにどのような行動機能が必要なのかを考えるための実務モデル**として用います。測定モデルとしての妥当性や信頼性については、今後さらに実証していく必要があります。

したがって、CAGE で見たいのは、「この人はどの種類の人間か」ではありません。チームが目的を実現するために、現在どの機能を必要としているのかです。

図 5-1

CAGE適応モデル

— 4つの適応スタイルが、個の力と環境のダイナミクスをつなぐ —

人はそれぞれ異なる強みで適応し、組織の中で互いに補完し合いながら成果を生み出す。CAGEは、変化と安定、自律と共生のバランスで人の可能性を可視化するモデルである。



CAGE は、「あなたは何タイプか」を決めるためのものではない

四象限のモデルを示すと、多くの人はすぐに、「私はどのタイプなのか」と考えます。しかし、CAGE 適応モデルで最も避けたいのが、この固定的な使い方です。

例えば、「あなたは Guardian です」と決めてしまえば、本人は「私は Guardian だから、新しいことを考えるのは得意ではありません」と説明できるようになります。診断結果が、成長しない理由になってしまうのです。

CAGE で問いたいのは、「あなたは何者か」ではありません。「今、この Role、この Context で、どの行動機能が必要なのか」です。

普段は安全(Guardian)の行動が多い人でも、新規事業では探究(Adventurer)を求められることがあります。成果(Challenger)を強く出す人でも、チームが疲弊している状況では関係(Empathizer)を意識する必要があるま

す。反対に、人間関係を重視してきた人でも、業績が大きく悪化している局面では成果(Challenger)として厳しい判断を行う必要があるでしょう。

必要な行動は、Role と Context の組み合わせによって変わります。人を固定するのではなく、状況に応じて必要な行動を選び直す。その意味で、CAGE は「適応モデル」なのです。

強みは、使いすぎれば弱みに変わる

CAGE の四つの機能に優劣はありません。成果も、探究も、安全も、関係も、組織には必要です。ただし、どの機能も使いすぎれば問題になります。

成果(Challenger)が強すぎれば、短期的な結果を優先するあまり、人を疲弊させたり、リスクを軽視したりすることがあります。探究(Adventurer)が強すぎれば、新しいアイデアは次々と出ても、いつまでも実行へ移らないかもしれません。安全(Guardian)が強すぎれば、品質は守られても、新しい挑戦が始まらなくなります。関係(Empathizer)が強すぎれば、雰囲気は良くても、必要な対立や厳しい指摘を避ける可能性があります。

したがって、CAGE で見るべきなのは、「どの機能が強いのか」だけではありません。**現在の状況に対して、どの機能が不足し、どの機能を使いすぎているのか**を見る必要があります。

ここで、第4章で扱ったメンタルモデル(Mental Models)ともつながります。「リーダーは必ず結果を出さなければならない」という前提を強く持つ人は、どのような状況でも成果(Challenger)を使おうとするかもしれません。しかし、現在の問題がチーム内の信頼関係の崩壊であれば、まず必要なのは関係(Empathizer)かもしれません。

成長とは、自分が得意な行動をさらに強くすることだけではありません。**必要なときに、自分が普段あまり使わない行動も選べるようになることです。**

目指すのは、「四種類の人」ではなく「四つの機能 を使えるチーム」

CAGE 適応モデルの目的は、Challenger、Adventurer、Guardian、Empathizer を一人ずつ集めれば良いチームになる、ということではありません。一人が複数の機能を担うこともありますし、同じ人でも状況によって使う機能は変わります。

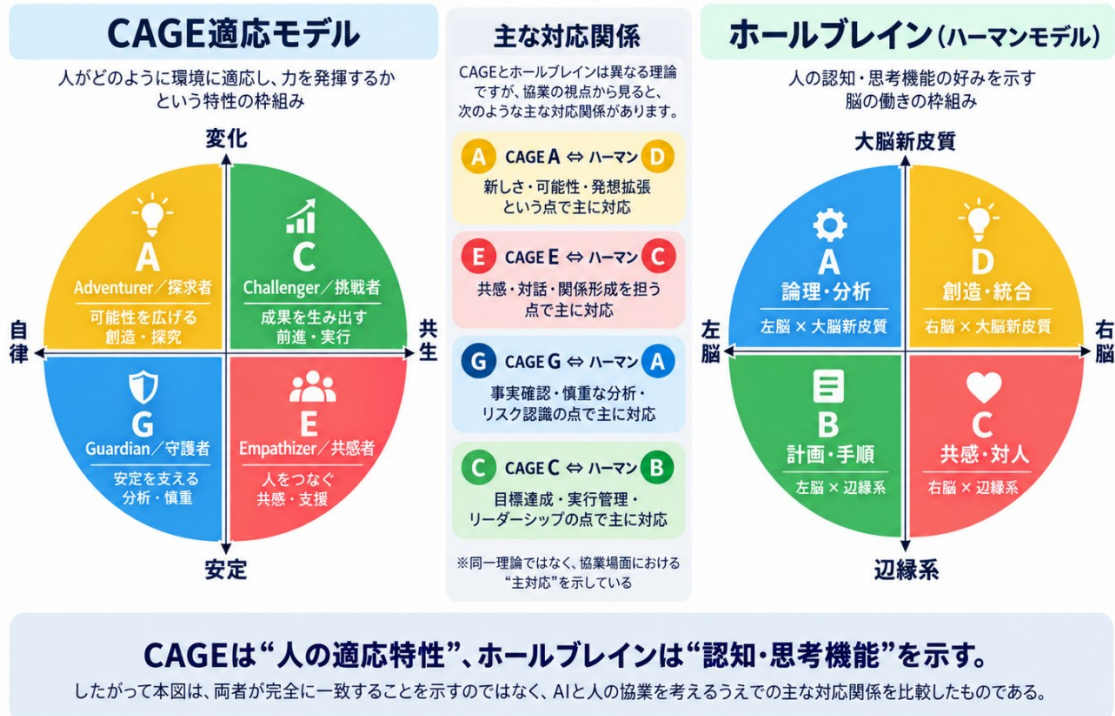
例えば、新規事業の初期には、新しい可能性を探す探究(Adventurer)が強く必要になります。しかし、アイデアだけでは事業にはなりません。ある段階からは、成果(Challenger)によって顧客を獲得し、数字へ変える必要があります。事業が拡大すれば、安全(Guardian)によって品質、手順、リスクを管理しなければなりません。人数が増え、部門間の連携が複雑になれば、関係(Empathizer)によって信頼と協働をつくる必要があります。

同じチームでも、時間とともに必要なバランスは変わります。重要なのは、チーム全体として、四つの機能を必要なときに使えることです。

私は、このように異なる思考・行動機能をチームとして補完し合える状態を、ホールブレイン的な状態（ハーマンモデル）と考えています。ただし、CAGE 適応モデルと既存の全脳思考(Whole Brain Thinking)そのものが同じ理論だという意味ではありません。世の中には、人や組織の違いを整理するさまざまなモデルがあります。CAGE の独自性を四象限という形そのものに求める必要はありません。

図5-2 CAGEとホールブレイン（ハーマンモデル）の比較

— 人の適応特性と認知機能の“主対応”を比較する —



本書でCAGEを使う意味は、四つの機能を **Role**、**Context**、**実際の行動**、そして**週次の学習**へつなげることにあります。

Shared Vision から Team CAGE へ

ここまでの流れを整理すると、Roleを決める順番が見えてきます。

最初に共有ビジョン(Shared Vision)があります。本人は何を目指すのか、会社は何を目指すのか、その接点はどこにあるのかを確認します。次に、会社の方向を「このチームは現在、何を実現しなければならないのか」というチーム目標へ落としします。そのうえで、その目標を実現するために、成果、探究、安全、関係のどの機能が必要なのかを考えます。これがTeam CAGEです。

例えば、新しい市場へ進出するチームであれば、新市場について学び、新しい顧客ニーズを発見する探究(Adventurer)が必要です。しかし、調査だけで終わらず、実際の受注へ進める成果(Challenger)も必要になります。未知の市場へ進む以上、契約、資金回収、品質、法規制などを見る安全(Guardian)も欠か

せません。そして、本社、現地、営業、管理部門など異なる人たちをつなぐ関係(Empathizer)も必要です。

そこで初めて、「では、このチームの中で、誰に何を担ってもらうのか」を考えます。

流れは、**共有ビジョン(Shared Vision)→チーム目標→必要な CAGE 機能→Role 設定**となります。

これは、「Aさんはこのタイプだから、この仕事を任せる」という発想とは逆です。先に実現すべき未来を見る。次に、その未来に必要なチーム機能を見る。そして最後に、人と Role を組み合わせる。これが、本書で考える CAGE を使ったチーム設計です。

Role は、肩書ではない

本書でいう「役割」(Role)は、営業部長、マネージャー、経理担当、コンサルタントといった肩書や職種そのものではありません。同じ営業部長でも、置かれているチームの状態によって期待されることは異なります。

停滞しているチームなら、新しい可能性を探る Role が重要かもしれません。急成長して混乱しているチームなら、安全をつくる Role が重要になるでしょう。人間関係が崩れているチームなら、関係を修復することが優先されます。

つまり Role とは、「このチーム、この目標、この Context の中で、あなたに何を担ってほしいのか」という期待です。チームや状況が変われば、その期待も変わります。

だから Role は、人に貼りつけられた固定ラベルではありません。チームの現在地に応じて、継続的に見直されるものです。

上司の仕事は、「人を評価する」だけでなく 「チームを設計する」こと

このように考えると、上司の仕事も変わります。

通常の人事評価では、「Aさんは成果を出したか」「Bさんは期待どおりだったか」と、一人ひとりを縦に見ることが中心になります。もちろん、個人評価は必要です。しかし、個人だけを見ていると、チームとして欠けているものが見えないことがあります。

例えば、営業チーム全員が今期の売上目標を達成しているとします。それでも、誰も新しい市場を探していなければ、半年後には成長が止まるかもしれません。開発チーム全員が納期を守っていても、誰も重要なリスクを指摘しないのであれば、大きな事故につながる可能性があります。個人としては高く評価できても、チームとして必要な機能が欠けていることがあるのです。

上司は、「誰が優秀か」だけでなく、「このチームは現在、何ができていて、何ができていないのか」を見る必要があります。どの機能が十分に使われ、どの機能が不足しているのか。その不足は本人の能力の問題なのか、Role 設計の問題なのか、協働関係の問題なのか、評価制度や会議の進め方の問題なのか。そこまで視野を広げます。

上司の仕事は、人を評価することだけではありません。チームを設計し、チームが自分たちで学習できる状態をつくることでもあるのです。

Team CAGE Map は、「自己診断の平均」ではない

では、チームの CAGE 状態をどのように見ればよいのでしょうか。ここで、個人の自己診断を単純に平均するだけでは不十分です。

本人が「私は Adventurer です」と考えていても、実際の仕事では新しい問いをほとんど出していないかもしれません。反対に、自分では慎重な人間だと思っている人が、あるプロジェクトでは誰よりも新しい方法を試していることもあります。第4章で述べたように、自己認識と実際の行動は一致するとは限りません。

そのため Team CAGE Map で本当に見たいのは、性格診断の構成比ではありません。チーム目標、設定された Role、週次で実際に取られた行動、本人の振り返りなどから、**今、このチームではどの行動機能が実際に使われているのか**を見る必要があります。

例えば、「成果へ向かう行動は多いが、安全確認に関する行動が少ない」「関係を維持する行動は多いが、既存の前提を疑う探究行動が少ない」といった形です。AIは、こうしたパターンを整理し、チームへ鏡として返すことができます。

ただし、AIが「このチームは Guardian が弱い」と決めつけるのではありません。「現在の行動から見ると、安全機能が十分に使われていない可能性があります」と仮説を返し、その意味を最終的に判断するのは人間です。

個人への Feedback だけでは、チーム学習(Team Learning)にならない

Team CAGE Map でチームの状態が見えても、それを一人ひとりへ個別に返すだけでは、まだ個人学習です。「Aさんはもっと安全を意識してください」「Bさんはもっと探究してください」と個人へ伝えても、チームそのものは学習対象になっていません。

チーム学習(Team Learning)にするには、チーム自身が自分たちの状態を見る必要があります。

例えば上司が、「現在の私たちは成果へ向かう行動は強い。一方で、立ち止まってリスクを確認する行動が少なくなっています」とチーム全体へ返し、「では来週、チームとして一つだけ安全に関する行動を増やすとしたら、何を変えるか」を考えます。

ここで必ずしも「安全担当者を一人決める」必要はありません。会議に確認項目を入れる、重要案件では別の人が見直す、決裁前にリスクを一度言語化する、といったように、チームの仕事の仕方そのものを変えることもできます。

個人を直すのではなく、チームの働き方そのものを学習対象にする。ここで初めて、チーム自身が学習主体になります。

Team Feedback は、「誰が悪い」ではなく「何を変えるか」を問う

チームへのフィードバック(Team Feedback)でも、個人責任へ戻してしまえば意味がありません。「安全機能が弱いから、Aさんがもっと Guardian になってください」とするだけでは、「この問題はAさんの責任だ」という構図へ戻ってしまいます。

しかし、本当にAさん個人の問題なのでしょうか。納期が厳しすぎて、確認する時間そのものがないのかもしれませんが。上司がスピードだけを評価しているのかもしれませんが。リスクを指摘すると「慎重すぎる」と否定される文化なのかもしれません。そもそも安全確認が、誰の Role にも定義されていない可能性もあります。

ここまで見ると、CAGE はシステム思考(Systems Thinking)ともつながります。個人の行動だけでなく、チームの偏り、Role、協働関係、評価、会議、業務設計まで見るからです。

したがって、Team Feedback で問うべきなのは、「誰が足りないのか」ではなく、「チームとして何を変えるのか」です。

良いチームとは、対立のないチームではない

CAGE の四つの機能を考えると、良いチームについての見方も変わります。良いチームとは、全員が仲良く、いつも同じ意見を持っているチームではありません。

成果(Challenger)は、「議論は分かった。では、誰がいつまでに実行するのか」と前へ進めます。探究(Adventurer)は、「そもそも、別の方法があるのではないか」と既存の前提を疑います。安全(Guardian)は、「その方法には、このリスクがあります」と立ち止まらせます。そして関係(Empathizer)は、「この進め方で、メンバー同士が本当に協働できるのか」と人と人の関係を見ます。

この四つが本当に働けば、一定の緊張や対立が生まれるのは自然です。それ自体は悪いことではありません。

問題は、機能の違いが人格への評価に変わることです。「あなたは慎重すぎる」「あなたは無謀だ」「あなたは数字しか見ない」「あなたは人間関係ばかり気にする」となれば、チームは分裂します。

そうではなく、「安全(Guardian)の視点から見ると、何が心配ですか」「探究(Adventurer)の視点からこの前提を疑うと、何が見えますか」と、個人の性格ではなく機能として扱う。そうすれば、違いを相手の欠点ではなく、チームが必要とする異なる視点として使うことができます。

違いがあることが問題なのではありません。違いをチームの判断能力として使えないことが問題なのです。

多様性(Diversity)を、組織能力へ変える

企業では、多様性(Diversity)の重要性が語られています。異なる経験や背景を持つ人が集まることには、大きな意味があります。しかし、異なる人がいることと、その違いを仕事の中で使えることは同じではありません。

違う意見を言えば否定される。上司へ合わせる方が安全である。評価される行動が一つしかない。そのような組織では、どれほど多様な人を採用しても、実際の行動は次第に同質化していきます。

CAGE で見たい多様性は、属性や人数構成ではありません。異なる行動機能が、実際の仕事の中で使われているかを見ることです。探究の視点から既存の前提を疑える。成果の視点から決断し、前へ進められる。安全の視点から必要なときに止められる。関係の視点から分断した人たちを再びつなげられる。

違う人がいるだけでは足りません。違いが仕事の中で使われて初めて、多様性は組織能力になります。

Role は、一度決めて終わりではない

チームの状況が変われば、必要な Role も変わります。新規事業の初期には探究を期待されていた人が、事業化の段階では成果を担うようになるかもしれません。急成長して混乱しているチームでは、安全の Role を強める必要があります。メンバー間の関係が悪化すれば、関係機能を意識的に増やす必要がありますでしょう。

したがって、共有ビジョン(Shared Vision)から必要な CAGE 機能を考え、Role を設定したら終わりではありません。実際に行動し、その結果をチームへ返し、Role、協働関係、仕事の進め方を見直し、再び行動します。

チーム学習の循環は、次のようになります。

共有ビジョン(Shared Vision)→チーム目標→必要な CAGE 機能→Role→行動→Team Feedback→Role・協働関係・Team Action の調整→再行動

つまり、Role そのものも学習対象なのです。

AI は、チーム設計の支援者であって責任者ではない

CAGE と AI を組み合わせれば、「AI が最適なチーム編成を決めてくれるのではないか」と考える人もいるでしょう。確かに、AI には支援できることがあります。

現在のチーム行動を整理する。不足している可能性のある機能を示す。複数人に共通する行動パターンを見つける。設定した Role と実際の行動とのずれを示す。時間とともにチームの行動がどう変化したかを見る。こうした分析には大きな可能性があります。

しかし、誰へ何を任せるのか、誰をどこまで挑戦させるのか、本人の将来希望と会社の必要性をどう調整するのか、誰と誰を組ませるのか、どのリスクを取るのかという判断には、多くの状況 (Context) と価値判断が含まれます。

だから、AI はチーム設計(Team Design)を支援する。しかし、チーム設計の責任者にはならない。この章でも、Human in the Loop という原則は変わりません。

共有ビジョン(Shared Vision)は、Role を通じて仕事になる

ここまでの流れをつなげると、個人の未来と会社の未来が、どのように日常の仕事へ降りてくるのかが見えてきます。

本人には、自分が目指す未来があります。会社には、会社が目指す未来があります。その接点を見つけ、会社の方向をチームとして実現すべき目標へ落とします。次に、そのチーム目標に必要な CAGE 機能を考え、一人ひとりの Role を決めます。

つまり、個人の未来と会社の未来の接点が共有ビジョン(Shared Vision)となり、共有ビジョンがチーム目標へ落ち、チーム目標が Team CAGE によって必要機能へ翻訳され、その必要機能が Role へ落ちるという流れです。

ここまで来て初めて、会社のビジョンが本人の日常業務へつながります。

共有ビジョンは、「共感しました」で終わるものではありません。Role へ翻訳されて初めて、仕事になるのです。

チーム学習(Team Learning)とは、自分たちの働き方を学習対象にすることである

チーム学習(Team Learning)というと、チーム全員で研修を受けたり、対話をしたりすることを想像するかもしれませんが、それらも大切ですが、本書で考えるチーム学習は、それより広いものです。

チームとして目標を持つ。その目標に必要な機能を考える。Role を設定する。そして実際に行動し、その結果を見る。自分たちの偏りや不足を確認し、Role や協働関係、仕事の進め方を変え、また行動する。

この循環そのものがチーム学習です。一人ひとりの能力を高めるだけではありません。誰が何を担うのか、誰と誰がどう協働するのか、どの機能が不足し

ているのか、どの機能を使いすぎているのかを、チーム自身が継続的に問い直します。

チーム学習とは、チームが自分たちの働き方そのものを学習対象にすることです。

ここまで来て初めて、個人の学習がチームの学習へ変わります。

Role を決めただけでは、人は動かない

しかし、まだ一つ問題が残っています。

上司が、「今のチームには探究が足りないので、あなたには探究の Role を期待します」と伝えたとします。本人も、その意味を理解し、納得した。しかし翌朝、実際に仕事を始めると、「では、今日は具体的に何をすればよいのか」が分かりません。

「もっと主体的に」「もっと顧客志向に」「もっと挑戦して」「もっと周囲と協力して」という言葉も同じです。方向としては意味がありますが、そのままでは具体的な行動にはなりません。

Role は、まだ抽象的なのです。

次に必要なのは、Role を実際に行動できる大きさまで小さくすることです。「探究してください」ではなく、「今週、一人の顧客に、これまで聞いたことのない質問を一つするとしたら、何を聞きますか」というところまで落としします。

しかも、上司や AI がその行動を命令するものではありません。問いを返し、選択肢を考え、最後は本人が選びます。つまり、Role を期待行動へ変え、さらに一週間で実行できる行動へ変換する必要があります。

ここで必要になるのが、ナッジ(NUDGE)と週次行動(Weekly Action)です。

人の価値観を直接変えようとしない。まず一つの行動を変えてみる。その行動から経験が生まれ、経験を振り返ることで、自分自身の見方が変わることがあります。

次章では、Role を小さな行動へ変え、人を直接変えようとせずに行動変容と学習を起こす、ナッジ(NUDGE)と週次行動(Weekly Action)について考えていきます。

第6章 人を直接変えずに、行動を変える

——ナッジ(NUDGE)と週次行動(Weekly Action)

第5章では、本人が目指す未来と会社が目指す未来との接点を共有ビジョン(Shared Vision)として捉え、それをチーム目標へ落とし、CAGE 適応モデル(CAGE Adaptation Model)によってチームに必要な機能を考え、一人ひとりの役割(Role)へつなぐところまで見てきました。

しかし、Role を決めただけでは、人の行動は変わりません。

上司が「もっと主体的に動いてほしい」「もっと顧客志向になってほしい」「部下をもっと育ててほしい」「もっと新しいことへ挑戦してほしい」と期待を伝えることはできます。しかし、本人が翌朝仕事を始めたとき、具体的に何をすればよいのか分からなければ、その期待は行動にはなりません。

これは、私自身が長い間、人材育成で直面してきた問題でもあります。以前の私は、「考え方が変われば、行動が変わる」と強く考えていました。そのため、本人の考え方や価値観へ直接働きかけようとしていました。しかし、プロローグで述べたように、それを強くやりすぎれば、仕事の指導が本人の価値観への介入へ近づいてしまいます。

現在の私は、もう一つ別の経路を重視しています。考え方を直接変えようとする前に、まず小さな行動を一つ変えてみる。行動が変われば経験が変わり、その経験を振り返る材料が生まれます。そして本人が自分自身で何かに気づけば、結果としてももの見方まで変わることがあります。

つまり、人の成長には「考え方→行動」という流れだけでなく、「**行動→経験→内省→考え方**」という逆方向の学習経路もあります。

この循環を日常業務の中で回すために、EvaliA ではナッジ(NUDGE)の考え方と週次行動(Weekly Action)を組み合わせます。

「主体性を持て」は、行動ではない

人材育成では、主体性、顧客志向、責任感、リーダーシップ、挑戦、協働といった抽象的な言葉が数多く使われます。どれも重要ですが、そのままでは具体的な行動ではありません。

例えば上司が部下へ、「もっと主体性を持ってください」と伝えたとします。上司の頭の中には、「指示を待たずに動いてほしい」「自分の意見を持ってほしい」「自分で問題を発見してほしい」といった具体的な期待があるのでしょうか。

しかし、部下が考えている「主体性」と同じとは限りません。本人は、「勝手に判断してはいけないと思っていました」と考えているかもしれません。過去に自分で判断したとき、「なぜ相談しなかった」と叱られた経験がある可能性もあります。

そう考えると、「主体性がない」という一言の中には、本人の能力だけでなく、上司との関係、権限、Role、過去の経験、チームの文化など、さまざまな要因が含まれています。それにもかかわらず、「もっと主体的になりなさい」と繰り返しても、本人には何を変えればよいのか分かりません。

そこで必要なのは、抽象的な期待を、実際に観測できる行動へ翻訳することです。

Role を、期待行動へ翻訳する

第5章では、Role を「このチーム、この目標、この状況の中で、あなたに何を担ってほしいのか」という期待として考えました。しかし、Role もまだ抽象的です。

例えば、「若手社員の育成を担う」という Role を与えられた管理職がいるとします。それだけでは、本人が具体的に何をすればよいのかは決まりません。部下へ答えを教えるのか、質問によって考えさせるのか、一部の判断を任せるのか、仕事の後にフィードバック(Feedback)するのか、本人の将来について話

すのか。Role を実際の仕事へ落とすには、「その Role を果たしている人は、具体的にどのような行動をしているのか」まで考える必要があります。

そこで Role から、観測できる「期待行動」へ翻訳します。「部下育成を担う」という Role なら、「部下の考えを先に聞く」「一部の判断を任せる」「仕事が終わった後に振り返る」「自分の答えを伝える前に質問する」といった期待行動が考えられます。

ここで初めて、抽象的だった Role が日常の仕事へ近づいてきます。

期待行動を、一週間で試せる大きさまで小さくする

しかし、「部下へ任せる」という期待行動でも、実際にはまだ大きすぎる場合があります。誰に任せるのか、どの仕事を任せるのか、どこまで判断を渡すのか、失敗した場合にはどの段階で介入するのかが決まっていないからです。

人が行動できないとき、すぐに「意志が弱い」「やる気がない」と考えるべきではありません。行動そのものが大きすぎて、何から始めればよいのか分からないことがあるからです。

「英語を勉強する」「顧客との関係を深める」「部下を育成する」「新規顧客を増やす」といった言葉は目標ではあっても、今日何をするのかまでは示していません。

そこで EvaliA では、期待行動をさらに、一週間で実行できる大きさまで小さくします。例えば、「今週、部下一人に、一つの判断を任せる」「次の顧客面談で、説明を始める前に、相手が実現したいことを一つ質問する」「今週の会議で、現在のやり方について一つだけ『なぜこの方法なのか』と問いかける」といった形です。

こうなると、本人は具体的に実行するかどうかを選ぶことができます。

したがって、行動へ落とす流れは、**役割(Role)→期待行動→週次行動(Weekly Action)**という三段階になります。

ナッジ(NUDGE)は、「小さな行動」の別名ではない

ここで、ナッジ(NUDGE)について整理しておきます。

行動経済学におけるナッジとは、強制、金銭、罰則を用いず、人間の心理的特性（認知バイアス）を利用して、人々が自発的に望ましい行動を選択するように「そっとつつく（Nudge）」アプローチをいいます。つまり、本人の選択肢を禁止したり、命令したりするのではなく、選択の環境を工夫することによって、ある行動を取りやすくする考え方です。

したがって、「小さな週次行動を設定すること」そのものをナッジと呼ぶのは正確ではありません。本書でナッジの考え方から取り入れているのは、本人に正しい行動を命令するのではなく、AIによって**本人が自分で望ましい行動を選びやすい条件をつくる**ことです。

例えばAIが、「もっと顧客志向になってください」と指示する代わりに、「次の顧客面談で、相手をより理解するために一つだけ質問するとしたら、何を聞きますか」と問いかける。本人自身が考え、必要であればAIがいくつかの候補を整理します。それでも、最後にどれを実行するかを決めるのは本人です。

つまり、**ナッジ(NUDGE)は選択しやすい環境を設計する原則であり、週次行動(Weekly Action)は、その中で本人が選んだ具体的な学習実験だと整理すると分かりやすくなります。**

「本人が選ぶ」と「何でも自由」は違う

ここで、一つ重要な境界があります。本人が週次行動(Weekly Action)を選ぶといっても、会社の仕事をすべて本人の自由意思へ委ねるという意味ではありません。

法令を守ること、安全基準を守ること、情報セキュリティを守ること、会社として決められた必要な手順を守ること。こうしたものは、「やるかどうかを本人が選ぶ」対象ではありません。個人情報適切に扱うかどうか、安全確認を行うかどうか、法令に定められた手続きを守るかどうかを、学習実験として本人に選択させることはできません。

会社には、会社として定めるべき境界があります。その境界の内側で、「自分が Role を果たすために、今週どの行動を試すのか」を本人が選ぶのです。

つまり、**遵守すべきものは会社が決め、成長のために試す行動には本人の選択を入れる**という区別が必要です。これは、Human in the Loop を運用するうえでも重要な境界です。

なぜ一週間なのか

EvaliA では、次の行動を原則として一週間単位で考えます。ただし、これは「人間の成長には七日間が科学的に最適である」という意味ではありません。一週間は、あくまで実務上の運転周期です。

一か月後まで振り返らなければ、その間に仕事の状況が変わり、本人自身も何を試そうとしていたのかを忘れる可能性があります。一方、毎日確認すれば、本人にも上司にも負担が大きくなり、学習支援よりも管理や監視へ近づく危険があります。

一週間程度であれば、実際の仕事の中で一つの行動を試す機会があり、その経験がまだ具体的なうちに振り返ることができます。何をしたのか、何が起きたのか、なぜそうなったと思うのか、何を学んだのか、次はどうするのか。こうした振り返りを仕事から切り離さずに繰り返すためのリズムとして、週次という単位を使っています。

重要なのは一週間という数字そのものではありません。**行動と内省の間を空けすぎず、経験を学習へ変える循環を止めない**ことです。

大きな目標を、「小さな実験」へ変える

週次行動(Weekly Action)には、もう一つ重要な意味があります。それは、行動を「決意」ではなく「実験」として扱えることです。

例えば、「私は部下に仕事を任せられる管理職になります」と宣言すると、大きな自己変革になります。実行できなかったときには、「やはり自分にはできない」と、自分自身を評価してしまいやすくなります。

しかし、「今週、一つの案件について、部下に自分の案を先に出してもらおう」という行動なら、人格を変える宣言ではありません。一つの実験です。実際にやってみて結果を見る。うまくいけば、なぜうまくいったのかを考える。うまくいかなければ、何が起きたのかを振り返り、次の週に方法を変えることができます。

このとき重要なのは、成功か失敗かだけではありません。その実験から何を学んだのかです。

新しい行動を試すたびに成功を要求される組織では、人は失敗しない行動しか選ばなくなります。学習する組織(Learning Organization)に必要なのは、小さく試し、結果から学び、次の行動を調整できる状態です。

行動が先に変わると、考え方が後から変わることがある

第4章では、メンタルモデル(Mental Models)について考えました。例えば、「部下へ任せれば失敗する」と考えている管理職がいるとします。その人へ「もっと部下を信頼しなさい」と説得しても、簡単には見方は変わりません。

しかし、一つの小さな仕事だけを部下へ任せてみることはできます。実際に任せてみると、部下が予想以上に自分で考えた。少し時間はかかったが、結果も出した。そこで管理職自身はその経験を振り返れば、「すべてを任せるのは不安だが、仕事を選べば任せられるかもしれない」という新しい見方が生まれることがあります。

ここでは、考え方が先に変わったわけではありません。**行動によって別の経験をつくり、その経験が、それまでのメンタルモデルを問い直す材料になったのです。**

この順番であれば、上司が本人の価値観を説得して変える必要はありません。本人自身が経験し、自分の見方を修正します。

人の「心」を外から直接変えるのではなく、本人が自分で学び直せる経験をつくる。ここに、ナッジ(NUDGE)の考え方と週次行動(Weekly Action)を組み合わせる意味があります。

CAGE も、行動へ落ちて初めて意味を持つ

第5章で扱った CAGE 適応モデル(CAGE Adaptation Model)も、診断や概念の段階で止まれば、実際の行動は変わりません。

例えば、成果(Challenger)の機能が必要だと考えたとしても、「もっと成果を意識してください」だけでは具体的な行動になりません。止まっている案件を一つ選び、次の期限と担当を明確にする、といった形まで落とす必要があります。

探究(Adventurer)なら、「今週、一つだけ『なぜこのやり方なのか』と問いかける」「顧客一人に、これまで聞いたことのない質問をする」といった行動が考えられます。安全(Guardian)なら、重要案件について「もし失敗するとしたら、何が原因になるか」を事前に一度考える。関係(Empathizer)なら、最近十分に話せていないメンバー一人と、仕事上困っていることについて話してみる、といった行動になるかもしれません。

ただし、これらは全員に同じように適用する標準行動ではありません。本人の Role と Context に応じて、「**今週の自分にとって、どの行動を試すことに意味があるのか**」を本人自身が考えることが重要です。

CAGE が性格診断で終わらず、Role と週次行動(Weekly Action)へつながって初めて、実際のチーム学習(Team Learning)へ使えるようになります。

Role と Weekly Action の間には、「本人の選択」が必要である

Role から週次行動(Weekly Action)へ移る間には、必ず一つ入れるべきものがあります。それが、本人の選択です。

会社や上司が Role を決め、そのまま人工知能(AI)が具体的な行動まで決めてしまえば、社員にとっては命令が細かくなっただけです。それでは、自律的な学習とは言えません。

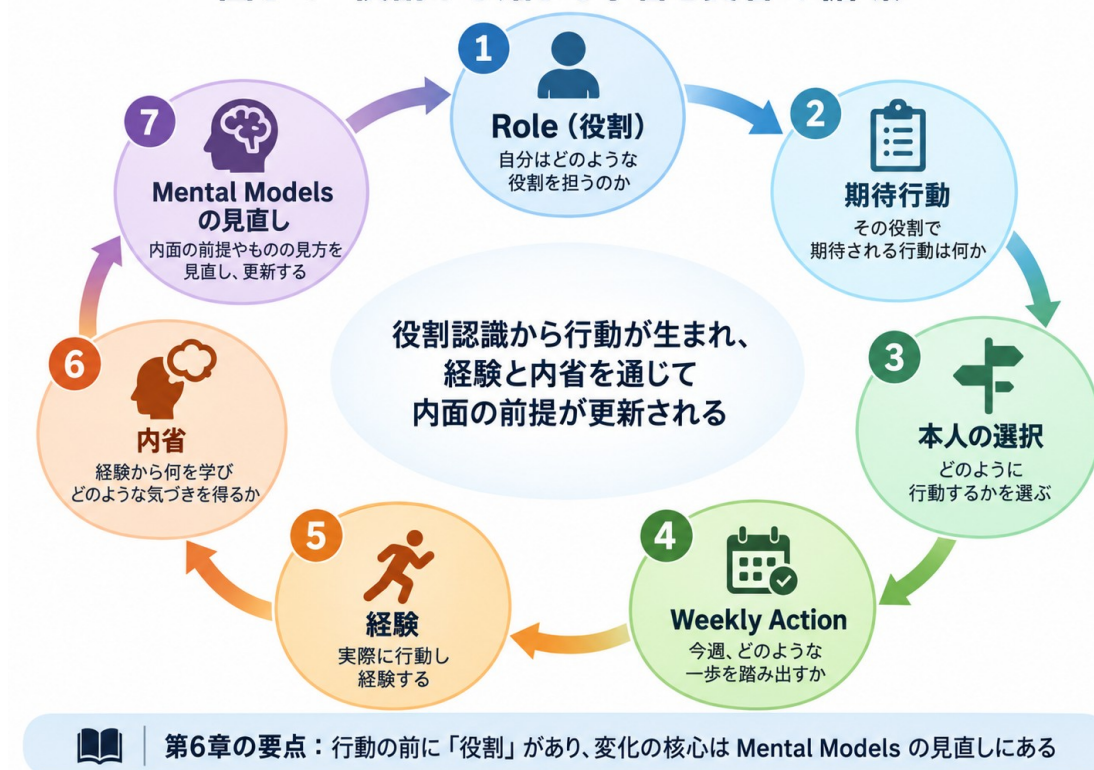
例えば AI が、「あなたには部下育成という Role が期待されています。今週試せる行動として、部下へ一つ判断を任せる、本人の話を最後まで聞く、仕事の後に短い振り返りを行う、といった選択肢があります。今の自分にはどれが必要だと思いますか」と問いかけることはできます。

本人は、その中から一つを選んでもよいし、自分で別の行動を考えてもかまいません。

したがって全体の流れは、**共有ビジョン(Shared Vision)→チーム目標→CAGE→Role→期待行動→本人の選択→週次行動(Weekly Action)→実行**となります。

この「本人の選択」が入ることで、会社が求める方向と本人の自律性を両立させることができます。

図6-1 役割から始まる学習と変容の循環



実行できなかったことも、重要な観測データになる

週次行動(Weekly Action)を決めても、実行できない週は当然あります。このとき、「約束したのに、なぜやらなかったのか」「本当にやる気があるのか」と本人の意思や人格へ原因を戻してしまえば、従来の管理に戻ってしまいます。

学習する組織では、未実行そのものを観測データとして扱います。

例えば、「次の顧客面談で、新しい質問を一つする」と決めたのに実行できなかったとします。単に忘れたのかもしれませんが。適切なタイミングがなかったのかもしれませんが。顧客に嫌われることを恐れたのかもしれませんが。過去に上司から「余計なことを聞くな」と注意された経験があるのかもしれませんが。そもそも設定した週次行動が、本人の Role や実際の仕事に合っていなかった可能性もあります。

同じ「未実行」でも、原因によって次に変える場所は異なります。本人の行動を変えるのか、メンタルモデル(Mental Models)を振り返るのか、Role を見直すのか、上司の支援を変えるのか、チームの働き方を変えるのか、仕事の仕組みそのものを変えるのか。

ここで重要なのは、本人のやる気だけを原因にしないことです。**できなかったという事実は、本人を評価する材料である前に、次にどこを変えるべきかを考えるための観測データです。**

未達でも、学習としては前進していることがある

週次行動が実行できなかったからといって、学習まで失敗したとは限りません。

例えば、「顧客へ提案する前に、相手の目的を聞く」と決めたのに、実際の商談ではできなかったとします。しかし振り返ってみると、本人は「沈黙が怖くて、相手が話し始める前に自分から説明してしまう」という自分の行動パターンに初めて気づいたかもしれません。

行動目標だけで見れば未達です。しかし学習という観点から見れば、それまで本人自身にも見えていなかったメンタルモデルや行動パターンが見えたことになります。

反対に、行動そのものは達成できても、上司に言われたから機械的に実行しただけで、本人が何も学んでいない場合もあります。

この違いを考えると、週次行動(Weekly Action)を単純な達成率だけで管理することが危険だと分かります。**行動の達成と、学習の前進は同じものではありません。**

成功も、「できました」で終わらせない

失敗だけではなく、成功した行動も重要な学習材料です。

週次行動を実行して成果が出たとき、「できました。よかったですね」で終わらせてしまえば、成功はその人だけの経験として消えていきます。

なぜ今回はできたのでしょうか。どの行動が効果的だったのでしょうか。他の場面でも再現できるのでしょうか。その人だからできたのでしょうか。それとも、他のメンバーにも使える方法なのでしょうか。

こうして成功を振り返ると、本人の偶然の成功だったものが、再現可能な知識へ変わる可能性があります。さらに、それをチームへ共有すれば個人学習がチーム学習(Team Learning)へ変わり、仕事の標準や仕組みまで変えられれば組織学習になります。

学習する組織は、失敗からだけ学ぶわけではありません。**失敗は仕組みを修正する材料になり、成功は仕組みに新しい能力を加える材料になります。**

この考えは、後に扱う名人循環(Mastery Loop)と失敗防止循環(Foolproof Loop)へつながっていきます。

ナッジ(NUDGE)が、操作へ変わらないために

ナッジ(NUDGE)や AI を人材育成へ使う場合には、避けて通れない問題があります。

会社には、社員に一定の方向へ行動してほしいという意図があります。そのため、ナッジは使い方を誤れば、本人の学習を支援する仕組みではなく、会社に都合のよい行動へ本人を誘導する技術になる可能性があります。

その境界を守るために、本書では三つの原則を重視します。第一に、本人が何を目指すのかについて、本人自身の意思を入れることです。会社の目標だけから週次行動をつくりません。第二に、AI や上司が行動候補を示しても、遵守事項を除き、学習行動の最終的な選択は本人に残します。そして第三に、実行しなかったことを直ちに人格評価へ結びつけません。

この三つがなければ、「選択を支援する」という名目で、実際には社員の行動を細かく統制することができます。

目的は、人を会社の都合どおりに動かすことではありません。**本人が自分で選び、自分の経験から学び、次の行動を自分で変えられる状態をつくること**です。

週次行動(Weekly Action)は、評価項目ではなく学習実験である

EvaliA を運用するときには、もう一つ注意が必要です。毎週、週次行動(Weekly Action)を設定すると、会社はつい「何%実行したか」を評価したくなります。

しかし、週次行動そのものを人事評価項目にしてしまえば、本人は達成しやすい行動ばかりを選ぶようになる可能性があります。100%達成できる行動を選んでおけば、評価を下げずに済むからです。それでは、学習実験としての意味が失われます。

週次行動は、基本的には学習のための小さな実験です。一方、会社として達成すべき成果(Result)や重要業績評価指標(Key Performance Indicator)は別に存在します。この二つを混同しないことが重要です。

成果には責任を持つ。しかし、その成果へ近づく過程では、小さな実験を行い、失敗から学ぶ余地を残す。この二層構造がなければ、学習する組織ではなく、毎週の行動まで細かく管理する組織になってしまいます。

人を直接変えない。しかし、放置もしない

本人の価値観を直接変えないというと、「それなら本人の自由に任せればよい」と聞こえるかもしれませんが、しかし、本書で考えているのは放任ではありません。

会社は目指す方向を示します。上司は Role を明確にします。AI は、本人が行動を考えるための問いや選択肢を返します。本人は週次行動を選び、実際に行動します。そして翌週、その経験を振り返ります。

できなければ、本人を責める前に原因を考える。うまくいけば、なぜうまくいったのかを振り返り、次の行動へつなげる。つまり、**強制はしない。しかし、問い続ける。**

本人の自律性を尊重することと、成長への関心を失うことは、まったく違います。

人の「心」を変えるのではなく、学習する条件をつくる

ここまで来ると、私自身の人材育成に対する考え方が、以前とは大きく変わったことが分かります。

以前の私は、「本人の考え方をどう変えるか」を強く考えていました。現在は、「**本人が自分の考え方を見直したくなるような経験を、どうつくるか**」を考えています。

上司が正しい考え方を教えるのではありません。AI が正しい価値観を判定するのでもありません。本人が小さな行動を選び、実際にやってみる。その結果として何かが起こり、それを振り返る。その中で本人自身が気づき、必要なら次の行動を変える。この循環を継続します。

こう考えれば、人材育成は価値観の矯正ではなく、**本人が自分自身の経験から学べる条件をつくること**へ変わります。

行動が変わっただけでは、「成長した」とは言えない

しかし、ここで新しい問題が出てきます。

週次行動(Weekly Action)を実行し、実際の行動が変わったとしても、それだけで本人が成長したと判断してよいのでしょうか。

必ずしもそうではありません。上司に言われたから行動しただけかもしれません。行動はできていても、本人が強い負担を感じ続けているかもしれません。毎週努力しているのに、実際の成果にはつながっていない可能性もあります。

反対に、今週の行動は未達でも、その理由を振り返る中で本人が重要なことに気づき、次の行動が大きく変わることもあります。

つまり、人の成長を知るには、「何をしたか」という一つの情報だけでは足りません。行動はどのように進んでいるのか。本人はその状態をどのように受け止めているのか。そして、最終的にどのような成果につながっているのか。この三つを分けて見る必要があります。

そこで次章では、EvaliA が人の学習を観測する三つの窓、**行動過程(Process)・心理状態(Mind)・成果(Result)**について考えます。

三つを単純に足して一つの点数にするわけではありません。その間に生まれるGap から、本人、上司、チーム、組織のどこを変えるべきなのかを見ていきます。

第7章 学習をどう観測し、どこを変えるか

——行動過程(Process)・心理状態(Mind)・成果(Result)と Gap

第6章では、役割(Role)を具体的な期待行動へ翻訳し、さらに本人が一週間で試すことのできる週次行動(Weekly Action)まで小さくすることで、日常の仕事そのものを学習の場へ変える方法を考えました。

しかし、行動が変わったからといって、それだけで「成長した」と判断することはできません。毎週、新しい行動を試しているのに成果(Result)にはつながっていないかもしれません。反対に、成果は出ていても、その成果が偶然によるもので、次に再現できないこともあります。また、外から見れば順調に働いているようでも、本人自身は強い負担や違和感を抱えているかもしれません。

つまり、人の学習を一つの窓だけから見ることにはできません。

そこで EvaliA では、人の学習状態を、**行動過程(Process)**、**心理状態(Mind)**、**成果(Result)**という三つの異なる経路から観測します。そして、三つを最初から一つの点数へまとめるのではなく、それぞれをできるだけ別の経路から見た後で、その間にどのようなずれ(Gap)があるのかを比較します。

この Gap が、次に何を詳しく見るべきか、どこへ働きかけるべきかを考える入口になります。

成果(Result)だけでは、成長の途中は見えない

会社にとって、成果(Result)を見ることは当然です。売上、利益、納期、顧客獲得、案件完了、品質など、仕事には達成すべき結果があります。人材育成だからといって、成果を問わなくてよいわけではありません。

しかし、成果だけを見ると、そこへ至る途中が消えてしまいます。

例えば、ある営業担当者が大きな契約を獲得したとします。結果だけを見れば、大きな成果です。しかし、その契約が本人の新しい営業行動によって生まれたのか、それとも長年の既存顧客から偶然大型案件が入ったのかでは、意味が違います。

前者であれば次の成果へつながる学習が起きている可能性があります。後者であれば同じ結果が再現できるとは限りません。

反対に、目標を達成できなかった人が、何も成長していないとも限りません。新しい営業方法を試し、これまで接触できなかった顧客へ入り始めているかもしれない。今期の数字にはまだ表れていなくても、将来の成果につながる行動変化が起きている可能性があります。

成果(Result)は重要です。しかし、成果だけでは、「なぜその結果になったのか」「次も再現できるのか」「本人はどのように変化しているのか」までは分かりません。

だから、成果へ至る途中を見る必要があります。それが、行動過程(Process)です。

行動過程(Process) —— 本人は、どのように前へ進んでいるのか

行動過程(Process)は、主として毎週の AI コーチングから観測します。

前回、本人は何をすると決めたのか。その行動を実際に試したのか。前へ進んだのか、止まったのか。何が阻害要因になったのか。そこから何を学び、次にどの行動を選んだのか。こうした一連の変化を見ていきます。

ここで見たいのは、単純な「実行率」ではありません。

例えば、四週間のうち三週間、週次行動(Weekly Action)を実行したから75%で良好、と評価したいではありません。三週間で何を試し、何を経験し、そこから何を学んだのか。実行できなかった一週間には、何が起きていたのか。その原因は本人自身にあったのか、それとも Role、上司、チーム、仕事の仕組みなどにあったのかを見る必要があります。

したがって、行動過程(Process)とは、人を細かく管理するための行動履歴ではありません。**本人が何を試し、どこで止まり、何を学びながら前へ進んでいるのかを見る学習の履歴**です。

心理状態(Mind) —— 本人は、その経験をどう受け止めているのか

行動過程(Process)だけを見ても、人の状態を十分には理解できません。同じ行動を取っている二人が、同じようにその仕事を受け止めているとは限らないからです。

例えば、二人の社員が同じように難しい仕事へ挑戦しているとします。一人は苦勞しながらも、「自分は成長している」と感じている。もう一人は同じように働いていても、「この仕事を続ける意味が分からない」と感じているかもしれません。外から見える行動が似ていても、本人がその経験をどう受け止めているかは大きく違います。

そこで EvaliA では、毎週の行動確認とは別に、一定の節目で本人自身が振り返る対話を行います。何に満足しているのか。何に不満を感じているのか。何に達成感を持ち、何を成長と感じているのか。逆に何が負担なのか。現在の仕事、上司、会社、自分の将来をどのように捉えているのかを言葉にしていきます。

ただし、ここは非常に重要です。本書でいう心理状態(Mind)は、医学的・臨床的な精神状態を AI が診断するという意味ではありません。また、AI が本人の内面そのものを直接読むこともできません。

本書で扱う心理状態(Mind)とは、**仕事について本人自身が語る満足、不満、意味、負担、成長実感などの自己認識や心理的シグナルを見るための窓**です。本人の内面そのものを「測定した」と考えないことが重要です。

成果(Result) —— 上司は、何を成果として見ているのか

三つ目が成果(Result)です。

成果については、本人の自己認識だけでは十分ではありません。会社として重要業績評価指標(Key Performance Indicator)や目標を設定し、本人へ Role を期待している以上、上司から見た成果確認も必要です。

ただし、ここでも単純な点数評価だけで終わらせません。どの目標が達成されたのか。本人のどの行動や判断が成果へ寄与したのか。未達だった場合、その理由は何か。本人がコントロールできた要因と、上司、組織、市場など本人だけではコントロールできなかった要因をどう分けるのか。今回の成果は、次回も再現可能なのか。こうしたことを確認します。

AI は、この確認のために上司へ問いを返すことができます。しかし、AI が「この社員は 82 点です」と最終評価を決めるわけではありません。評価責任を持つのは人間です。

AI は、上司が印象だけで判断しないように、事実を具体化し、別の可能性を問いつけるための支援をします。したがって、成果(Result)も、評価を確定するためだけの情報ではありません。**成果を確認し、その結果を次の学習へ戻すための観測窓**でもあります。

三つの窓を、最初から一つに混ぜない

EvaliA では、行動過程(Process)、心理状態(Mind)、成果(Result)を、できるだけ一方の評価に引っ張られない形で取得します。

例えば、上司が先に「この社員は成果を出している」と考えている状態で本人の心理状態を推測すれば、「本人も順調なはずだ」と見てしまう可能性があります。

反対に、本人が「自分は頑張っている」「成長を実感している」と話しているからといって、実際の成果(Result)まで良いとは限りません。また、本人が毎週行動しているからといって、その行動が正しい方向へ向かっているとも限りません。

そのため、三つを最初から一つの評価へまとめず、**できるだけ別々に取得し、後から比較することが重要になります。**

もちろん、三つが完全に独立しているという意味ではありません。同じ本人が、同じ仕事の中で経験していることですから、互いに関係しています。しかし、最初から一つの評価に混ぜてしまえば、「本人は順調そうだから成果も良いはずだ」「上司の評価が高いから本人も満足しているはずだ」といった思い込みが入り、重要な違いが見えにくくなります。だからこそ、先に混ぜるのではなく、後から比較するのです。

三つの間にある Gap を見る

三経路観測の目的は、Process、Mind、Result という三つのスコアをつくることではありません。その間にどのような Gap があるのかを見ることです。

特に見るのは、**Process と Mind**、**Process と Result**、**Mind と Result** の間にある違いです。

例えば、行動過程(Process)は前進しているのに、心理状態(Mind)では負担や不安が強まっていることがあります。この場合、「行動できているから問題ない」とは言えません。現在の Role が本人の目指す未来とずれているのかもしれない。上司からの支援が足りない可能性もあります。あるいは、新しい Role へ挑戦している途中で、一時的に負荷が高まっているだけかもしれません。

反対に、心理状態(Mind)は良好なのに、行動過程(Process)が長く停滞している場合もあります。本人は現在の仕事に満足していても、新しい学習が止まり始めている可能性があります。

Process が前進しているのに Result が出ない場合もあります。本人が十分に行動しているのであれば、対象顧客、価格、商品、KPI、上司の支援、業務プロセスなど、本人以外の構造を見る必要があるかもしれません。

逆に Result が良くても、Process を見ると、将来へつながる学習が起きていないこともあります。長年の既存顧客から大型案件が入ったため数字は良いが、新しい顧客をつくる行動は止まっている。ある管理職が高い業績を出しているが、すべての判断を自分で抱え込み、部下へ何も任せていない。短期的な成果は良くても、再現性や持続性には問題があるかもしれません。

Mind と Result の間にも同じことが起きます。成果は高いのに、本人が強い負担を感じている。反対に、本人は仕事に満足しているのに、成果が出ていない。

どれか一つだけを見れば、異なる結論になります。だから、違いそのものを見るのです。

図 7-1

人とチームの成長を支える3つの視点とGap

— Process・Mind・Resultを別々に観て、その違いから次の一步を見つける —



Gap は、「原因」ではなく「探索地点」である

ここで最も重要なのは、Gap 自体を問題や原因と決めつけないことです。

例えば、Process は前進している。Result も良い。しかし Mind では負担が強くなっている。この状態を見て、AI が「この社員は燃え尽きています」と断定してはいけません。

一時的に難しい仕事へ挑戦しているだけかもしれません。本人自身が短期間の負荷を受け入れている可能性もあります。Role と評価制度とのずれが不安を

生んでいるのかもしれませんが。あるいは、節目の対話で語られた内容だけでは、本人の状態を十分に表していない可能性もあります。

Gap が示すのは原因ではありません。「ここには、もう少し詳しく見る価値がある」という探索地点です。

AI が行うのは、原因を確定することではありません。「本人の成果認識と現在の Role の間にずれがある可能性はありませんか」「上司と本人で評価基準の理解が違っていませんか」「複数人に同じ現象が起きているなら、制度側も確認する必要があるのではありませんか」と、原因仮説を広げることです。

その仮説を本人や上司が実際の状況と照らし合わせ、必要ならさらに事実を確認します。最終的に、その Gap にどのような意味があるのかを考え、どこを変えるのかを決めるのは人間です。

Gap は原因ではなく、原因を探す場所なのです。

Gap を見つけたら、「誰を直すか」ではなく「どこを変えるか」

三つの観測から Gap が見つかり、次に必要になるのが「介入」(Intervention)です。

ここでいう介入とは、人の心を操作することではありません。学習を前へ進めるために、どこへ働きかけるべきなのかを決めることです。そして、変える対象は本人とは限りません。

本人自身で変えられる問題もあります。上司のフィードバック(Feedback)、支援、権限の渡し方を変える必要がある場合もあります。チームの Role や協働関係を見直すべき場合もあります。さらに、標準、制度、業務プロセス、システムなど、組織側を変えなければ解決できない問題もあります。

本書では、介入先を大きく **本人、上司、チーム、組織** という四つに分けて考えます。

この四層を持つことで、「問題が本人のところに現れているから、本人を教育する」という短絡を避けることができます。

本人で変えられるなら、次の Weekly Action へ戻す

本人自身で変えられる問題であれば、次の週次行動(Weekly Action)へ戻します。

例えば本人が、「顧客へ質問したいのに、つい自分から説明を始めてしまう」と気づいたとします。この場合、次の一週間に、「顧客との最初の五分間は、自分の説明より先に一つ質問する」という小さな実験を行うことができます。

実行し、その結果を振り返り、次の行動を変える。これは「個人学習」です。

すべての問題を大きなチーム課題や組織課題へ上げる必要はありません。本人自身で変えられるものは、できるだけ本人の学習循環へ返します。

上司の関わり方を変える必要があることもある

一方、本人だけでは解決できない問題もあります。

例えば本人は、「もっと仕事を任せてほしい」と考えている。しかし上司は、「まだ任せられない」と考えている。この場合、本人へ「もっと主体的に動いてください」と言っても問題は解決しません。

何を任せなのか。どこまで本人が判断できるのか。どこから上司へ相談するのか。失敗をどこまで許容するのか。こうした Role や権限の設計を明確にする必要があります。

また、上司の Feedback そのものが本人の学習を止めていることもあります。本人が何か新しいことを試すたびに、上司がすぐ正解を教えてしまえば、本人が自分で考える余地は小さくなります。

このような場合には、本人ではなく、上司の支援、Feedback、権限移譲、Role 設定を変えることが介入になります。

複数人に同じ Gap があれば、チームを見る

同じような問題が複数人に起きている場合には、一人ひとりの能力だけでは説明できない可能性が高まります。

例えば複数の社員が、「新しい提案をすると、失敗したときに責任を問われるので動きにくい」と感じているとします。その場合、一人ひとりへ「もっと挑戦してください」と促しても、根本的な問題は変わりません。

評価の仕方、会議の進め方、チームの文化、CAGE 適応モデル(CAGE Adaptation Model)で見た行動機能の偏り、Role の配置、協働関係などを見直す必要があります。

第5章で述べた Team Feedback も、ここで使われます。個人を責めるのではなく、「チームとして、なぜこの行動が出にくいのか」を考え、Role、協働関係、チームとしての行動を変える。これがチームへの介入です。

構造問題なら、SCAD へ上げる

さらに、本人や上司、チームだけでは変えられない問題があります。

同じ入力ミスが複数の部署で起きている。何人もの社員が同じ情報を二重入力している。誰もが同じ承認待ちで仕事を止めている。顧客から同じ不満が何度も出ている。

このような問題を、「次回から気をつけよう」で終わらせてはいけません。

ここで SCAD を起動します。

SCAD では、問題を共有(Share)し、本当に構造問題なのかを確認(Check)し、手順、標準、ルール、システムなどを調整(Adjust)し、変更した仕組みを実行(Do)します。そして、変更後に同じ問題が減ったのかを、再び EvaliA で観測します。

つまり、EvaliAで観測する→構造問題を見つける→SCADで仕組みを変え
る→新しい仕組みで実行する→再び EvaliA で観測するという循環になります。

ここで、個人の学習が組織の仕組みの学習へつながります。

共通 Gap は、組織が学ぶべき場所を教える

AIを使う大きな意味の一つは、一人ひとりを個別に見るだけでなく、複数人の間に共通するパターンを見つけやすくなることです。

例えば、一人の社員だけが「承認に時間がかかって仕事が進まない」と話しているなら、その案件固有の問題かもしれません。しかし、同じ部署の多くの社員が同じところで止まっているなら、承認プロセスそのものを見る必要があります。

また、複数の管理職が「部下が自分で判断しない」と感じている一方で、複数の部下が「自分で判断すると後から修正されるので、先に上司へ聞いた方が安全だ」と話しているとします。この二つを並べれば、「主体性のない社員が多い」という問題から、「自律的な判断を生みにくい組織構造になっているのではないか」という問いへ変わります。

このように、一人の中では個人課題に見えていた Gap が、複数人に共通して現れることで、チームや組織側の課題として見えてくることがあります。

ただし、共通 Gap が見えたからといって、「これは制度の問題だ」と AI が原因を決めるわけではありません。共通 Gap が示すのは、「**本人だけではなく、構造側も調べる必要がある**」ということです。

AI は、原因を決めるのではなく、探索範囲を広げる

Process、Mind、Result を比較し、さらに複数人の共通 Gap まで見られるようになると、AI へ原因判断まで任せたくなるかもしれません。

しかし、ここには慎重が必要です。AI ができるのは、利用を許可された情報を比較し、「この可能性も考えられませんか」と原因仮説を増やすことです。

本人の問題に見えるものを、上司から見直す。上司との関係に見えるものを、チームから見る。短期的な問題を時間軸を延ばして見る。複数人に共通するなら、組織構造を見る。

AIが一つの原因へ絞り込むことよりも、人間が見落としていた可能性を広げることには価値があります。AIは、**原因の裁判官ではありません。原因探索の視野を広げる支援者です。**

Mindが悪いから励ます、Resultが悪いから教育する、ではない

心理状態(Mind)を見るようになると、本人の満足が低い、仕事への意味を感じにくくなっている、といったシグナルが見えることがあります。そのとき、すぐに「もっと褒めよう」「もっと励まそう」と考えるのは早すぎます。

原因がRoleの不明確さ、過大な業務量、評価制度への不信、本人の将来と現在の仕事とのずれであるなら、励ますだけでは何も変わりません。Mindは対策そのものを教えてくれるものではなく、**「何かを見る必要がある」ことを知らせるシグナル**です。

成果(Result)についても同じです。KPIが未達だったから研修を増やす、という対応は分かりやすい。しかし、本人が必要な行動を十分に取っているなら、さらに教育を増やしても結果は変わらないかもしれません。KPI自体が適切でないのかもしれません。商品や価格の問題かもしれない。上司が十分な権限を与えていない可能性もあります。

したがって、「Mindが悪い→励ます」「Resultが悪い→本人を教育する」という短絡を避ける必要があります。

Process、Mind、Resultを並べる意味は、本人だけを見ることから離れ、どこを変えるべきなのかを考える視野を広げることにあります。

観測の目的は、人を点数化することではない

三つの観測経路をつくると、Process、Mind、Result にそれぞれ点数をつけ、人を精密に数値化したくなる誘惑があります。

しかし、仮に Process が 78 点、Mind が 62 点、Result が 81 点と表示されたとしても、それだけでは本人が何を学ぶべきなのか、会社がどこを変えるべきなのかは分かりません。

本書で重要なのは、三つを合計して一つの「人材スコア」をつくることではありません。どこが変化しているのか。どこに Gap があるのか。その Gap には、どのような理由が考えられるのか。次に何を確かめる必要があるのか。そして、誰が何を変えるのか。

観測の目的は、人を順位づけることではありません。次に学ぶべき場所と、次に変えるべき場所を見つけることです。

内省の情報と、組織学習の情報を分ける

心理状態(Mind)まで扱う以上、情報の使い方には明確な境界が必要です。

本人が AI との振り返りで、「今の上司には納得していない」「現在の会社方針には疑問がある」と話したとします。その生の対話がそのまま上司や評価者へ渡るのであれば、本人は次から本当に考えていることを話しにくくなるでしょう。

だから、本人が内省するための情報と、組織が学習するための情報を分けて考える必要があります。

本人との生の対話は、本人自身が振り返るための空間として扱う。一方、組織へ返す必要があるときには、複数人に共通する阻害要因、前進や停滞の傾向、共通 Gap、構造的問題のシグナルなど、組織学習に必要な形へ整理して返します。

誰が何を見ることができ、何を評価へ使い、何を本人の内省空間として扱うのか。その境界が曖昧なら、学習のための仕組みは簡単に監視へ変わります。

観測できることと、すべてを見てよいことは同じではありません。

三つの観測によって、人材成長循環(Human Development Loop)が閉じる

ここまでで、第2章で示した人材成長循環(Human Development Loop)が一周つながります。

計画(Plan)では、本人のありたい姿と会社との接点を考え、チーム目標、Role、重要業績評価指標(KPI)を設定します。実行(Do)では、Role を期待行動へ翻訳し、本人が週次行動(Weekly Action)を選び、実際の仕事の中で試みます。

確認(Check)では、行動過程(Process)、心理状態(Mind)、成果(Result)を、できるだけ別の経路から観測し、その間に生じている Gap を見ます。

改善(Act)では、Gap から原因を決めつけるのではなく、どこへ働きかけるべきかを考えます。本人の次の行動を変えるのか。上司の関わり方を変えるのか。Role やチームの協働関係を変えるのか。それとも SCAD によって仕事の仕組みそのものを変えるのか。そこを人間が判断します。

そして、その結果を再び次の Plan へ戻します。

つまり、**ありたい姿→Role・KPI→Weekly Action→Process・Mind・Result→Gap→Intervention→次の Plan** という循環です。

計画する。行動する。三つの窓から観測する。Gapを見つける。どこを変えるのかを決める。そして、もう一度試す。これが、本書でいう人材成長循環(Human Development Loop)です。

第二部を貫く原則 ——人を一つの情報だけで決めない

第二部では、人と組織が学ぶための設計を考えてきました。

第3章では、「心」と「技」という二つの方向から人の現在地を考えました。第4章では、本人が現在地を見るときに使っているメンタルモデル(Mental Models)を問い直しました。第5章では、共有ビジョン(Shared Vision)、CAGE 適応モデル(CAGE Adaptation Model)、Role を通じて、個人の学習をチームへつなぎました。第6章では、Weekly Action ができなかったという事実を、本人の意志の弱さだけで説明しないことを考えました。

そして本章では、Result だけで人を判断せず、Process、Mind との Gap から、次にどこを見るべきかを考えました。ここには、第二部全体を貫く一つの原則があります。

人を、一つの数字、一つの行動、一つの評価、一つの視点だけで決めない。複数の視点を持ち、その間にある違いから学ぶ。私は、この姿勢そのものが、学習する組織(Learning Organization)に必要なだと考えています。

学習する組織(Learning Organization)とは、ずれを消す組織ではない

ここまで見てくると、学習する組織について、もう一つ重要なことが分かります。

学習する組織とは、ずれのない組織ではありません。人と会社の方向はずれます。本人の自己認識と実際の行動もずれます。努力と成果もずれます。本人の受け止め方と上司の評価もずれます。チームに必要な機能と、実際に使われている機能もずれます。

問題は、ずれがあることではありません。ずれを見つけられないこと、ずれを隠すこと、そして、ずれを本人一人の責任にして終わることです。

学習する組織は、ずれを観測します。その意味を問い、原因を決めつけず、本人、上司、Role、チーム、仕組みのどこを変えるべきかを考えます。そして、もう一度試します。

AIは、その循環の中で正解を決める存在ではありません。問いを継続し、過去を記憶し、異なる情報を比較し、人間だけでは見落としやすいGapを見えるようにします。

そして、人間が意味づけし、選び、判断する。この関係によって、個人の学びをチームの学びへ、さらに組織の学びへ広げることができます。

では、これらの仕組みは実際の EvaliA の中で、どのように動くのでしょうか。最初に本人へ何を問い、毎週どのような対話を行い、本人、上司、チームへどのような情報を返していくのでしょうか。ここからは、設計思想から実際の運用へ進みます。

第三部 EvaliA は実際にどう動くのか

第8章 最初に「何になりたいか」を聞く

——本人の未来を、会社の未来と役割(Role)・重要業績評価指標(KPI)へつなぐ

第一部では、なぜ学習する組織(Learning Organization)を目指すのかを考え、第二部では、そのために人と組織の学習をどのように設計するのかを見てきました。

第三部では、一人の社員が EvaliA を使い始め、毎週の仕事を振り返り、チームとの関係を変え、一定の節目で自分自身の状態を見つめ、上司から成果を確認され、その経験が最後には会社の仕組みへ戻っていくところまでを追っていきます。

なお、ここで説明する機能には、2026年9月時点で実際に運用しているものと、設計・検証を進めているものが含まれます。終章で示す実証データは、運用中の機能から得られたものです。一方、将来の実装を含む部分については、EvaliA が目指す運転モデルとして記述します。

ここで登場する人物は、実在する特定の社員ではありません。私たちがこれまで人材育成や組織運営の中で経験してきた複数のケースをもとに構成した、説明のための架空のペルソナです。

佐藤健太、35歳。仕事のできる中堅社員

佐藤健太さんは35歳。海外事業を支援する会社で働いて10年になり、現在は数名のメンバーを持つチームリーダーです。

専門知識があり、顧客対応にも強い。難しい案件を任せれば、自分で調べ、自分で判断し、最後まで仕事を完結できます。上司からも顧客からも、「仕事のできる社員」として評価されてきました。

佐藤さんには、数年後に実現したいことがあります。海外拠点の責任者になることです。

単に海外で働きたいわけではありません。現地で顧客を増やし、現地社員を育て、自分の責任で一つの拠点を経営できる人になりたいと考えています。会社側も海外事業を拡大し、次世代の拠点責任者を育成する必要があるため、佐藤さんはその候補の一人です。

ところが、佐藤さんには一つ大きな課題があります。自分で仕事をするには非常に強いのですが、人へ仕事を任せることがあまり得意ではありません。

部下が困れば丁寧に教えます。仕事が遅れていれば自分が助けます。重要な顧客から質問が来れば、間違いがないように自分で答えをつくります。本人には、「自分は部下をよく支援している」という認識があります。

実際、佐藤さんがいると仕事は速く進みます。しかし、佐藤さんがいないと判断が止まりやすい。部下は困るたびに佐藤さんへ聞き、佐藤さんは答える。その繰り返しによって、佐藤さん自身の能力の高さが、結果としてチームの自律を遅らせている可能性があります。

このまま海外拠点責任者になれば、佐藤さん自身は誰よりも働くでしょう。顧客からも信頼されるかもしれません。しかし、あらゆる判断が佐藤さんへ集中する拠点になる可能性があります。

プレーヤーとしての強さを、そのまま管理職としての強さへ持ち込むだけでは、次の成長にはならないのです。

最初に聞くのは、会社の目標ではない

人材育成や評価の仕組みでは、最初に会社側の目標を設定することが多いでしょう。今期の売上目標はいくらか。重要業績評価指標(Key Performance Indicator)は何か。上司は何を期待しているのか。それらは、会社を運営するうえで当然必要な情報です。

しかし、EvaliA の計画(Plan)は、そこだけから始まりません。最初に本人へ聞きます。

「あなたは、何になりたいのですか。」

佐藤さんは、「海外拠点の責任者になりたいです」と答えます。

しかし、これだけではまだ肩書への希望にすぎません。そこで AI は、答えを与えるのではなく、本人自身の言葉を深めるための問いを返します。

なぜ海外拠点責任者になりたいのか。責任者になったとき、今の自分とは何が違ってほしいのか。責任者という肩書そのものが欲しいのか。それとも、責任者になることで実現したいことがあるのか。

こうした対話を続ける中で、佐藤さんは少しずつ自分の考えを言葉にしていきます。

「海外で、自分の責任で事業を動かせるようになりたいです。」

さらに、「事業を動かしている状態とは、具体的にはどのような状態ですか」と問われると、佐藤さんは考えます。

そして、「顧客から選ばれて、現地社員も育て、自分が細かいところまで見なくても仕事が回っている状態です。」と答えました。

ここで、「海外拠点責任者になりたい」という目標の意味が変わります。佐藤さんが本当に目指しているのは、単に役職に就くことではありません。**現地社員を育て、自分がいなくてもチームとして成果を出せる経営人材になること。**

EvaliA が最初に行うのは、本人へ会社の目標を押し込むことではありません。本人自身が、自分は何を実現したいのかを、自分の言葉で理解するところから始めます。

「何になりたいか」を、「何ができるようになりたいか」へ変える

将来の肩書が決まっただけでは、成長の計画にはなりません。次に考えるのは、その人になったときに「何ができるようになっていく必要があるのか」です。

佐藤さんには、すでに高い専門能力があります。顧客対応力もあり、一人で問題を解決する力もあります。

しかし、海外拠点責任者になれば、自分一人で仕事を終わらせることはできません。現地社員へ仕事を任せ、本人たちに判断させ、異なる能力を持つ人を組み合わせながら、チーム全体として成果を出す必要があります。

ここで、佐藤さんの現在地と目指す姿との間に、一つの大きな成長テーマが見えてきます。

「自分ができる」から、「人を通じてできる」へ。

ただし、AIがここで、「あなたの弱点は部下育成です」と診断するわけではありません。本人の言葉、これまでの経験、会社や上司から期待されていることを材料にしながら、現在地を見るための問いを増やします。

自己マスタリー(Personal Mastery)は、理想の未来だけを描くことではありません。ありたい姿と現在地の両方を見て、その間にある差を学習の対象にすることです。

本人の未来と、会社の未来を接続する

本人の未来が明確になっても、それだけでは会社の中の Plan にはなりません。次に、会社が目指している方向との接点を確認します。

佐藤さんは、海外拠点責任者になりたいと考えています。一方、会社も海外事業を拡大し、現地で事業を担える人材を増やしたいと考えています。ここには明確な接点があります。

すると、これまで「会社から与えられていた仕事」の意味も変わってきます。海外案件へ参加すること。現地社員と仕事をすること。部下へ仕事を任せること。拠点の数字を見ること。

それらが、会社のための業務であると同時に、佐藤さん自身が目指す未来へ近づくための経験になります。

本書で共有ビジョン(Shared Vision)を実務へ落とすときに重視しているのは、会社のビジョンを本人へ暗記させることではありません。**本人が目指す未来と会社が目指す未来の間に、本人自身が意味を感じられる接点を見つけること**です。

もちろん、いつも接点が見つかるとは限りません。本人は専門家として生きたいのに、会社は管理職を期待しているかもしれません。本人は国内で専門性を高めたいのに、会社は海外赴任を求めていることもあります。

その場合、AIが本人を会社の方向へ説得するものではありません。

まず、「現在はずれている」という事実を見えるようにします。そのうえで、Roleを変えれば接点をつくれるのか、別のキャリアがあるのか、本当に本人と会社の方向が異なるのかを、人間が考えます。

会社の戦略を、現在のチームの仕事まで下ろす

佐藤さんの未来と会社の海外戦略がつながっても、まだ日常の仕事にはなっていません。「海外事業を拡大する」という言葉では大きすぎるからです。そこで、現在のチームが何を実現しなければならないのかまで下ろします。

例えば、佐藤さんのチームに、「現地社員だけで、一定範囲の顧客対応を完結できる状態をつくる」という目標があるとします。この瞬間、本人の成長課題と会社の経営課題が、同じ仕事の中で重なります。

会社は、現地社員が自律して動ける状態を必要としている。佐藤さんも、将来海外拠点を経営するためには、人を通じて成果を出す力を身につける必要があります。

この二つを別々に考える必要はありません。

会社が必要としている仕事そのものを、本人の成長機会にすることができるからです。人材育成を、通常業務とは別の研修だけで行う必要はありません。日常の仕事そのものを、学習の場へ変えることができます。

そこで初めて、Role を置く

チーム目標が見えたところで、上司は「その中で、佐藤さんには何を担ってもらうのか」を考えます。

例えば、佐藤さんには次のような Role が設定されます。

「自分で顧客対応を完結するだけでなく、現地社員へ判断を移し、本人たちだけでも一定範囲の仕事を完結できる状態をつくる。」

ここでいう Role は、役職名ではありません。現在のチームの中で、本人に何を担ってほしいのかという期待です。そして、この Role には二つの意味があります。

会社にとっては、現地社員が自律して動けるチームをつくるために必要な Role です。一方、佐藤さんにとっては、将来、海外拠点責任者になるために必要な経験そのものです。

つまり、**会社の必要と本人の成長が、一つの Role で重なっています。**

ただし、Role を AI が自動的に決めるわけではありません。AI は、本人の希望、会社の方向、チーム目標などから、「こうした Role が接点になる可能性があります」と候補や問いを出すことはできます。

しかし、誰に何を任せるのかは、人間が判断します。Role には、責任だけでなく権限、本人の経験、チーム構成、顧客への影響など、多くの状況（Context）が含まれるからです。

Role と重要業績評価指標(KPI)をつなげる

Role が決まったら、その Role についてどの成果を確認するのかを考えます。

ここで重要なのは、Role と重要業績評価指標(Key Performance Indicator)を同じものにしないことです。

Role は、「何を担ってほしいのか」です。KPI は、「その Role について、どの成果を確認するのか」です。

佐藤さんの Role が「現地社員へ判断を移し、自律して仕事が進む状態をつくる」であれば、現地社員だけで完結した案件の割合、佐藤さん自身の直接介入なしで処理できる業務範囲、顧客対応品質などが KPI の候補になります。実際にどの KPI を使うかは、業務や戦略によって変わります。

ここで大切なのは、会社が期待する Role と、会社が評価する KPI を矛盾させないことです。

例えば会社が、「部下を育て、自律させてほしい」と佐藤さんへ求めながら、実際には佐藤さん本人の売上だけを評価していたらどうなるでしょうか。

短期的に自分の売上を最大化するためには、佐藤さん自身が仕事を抱えた方が合理的です。部下へ任せれば、一時的には時間もかかります。失敗する可能性もあります。

つまり、会社が口では「育成してください」と言いながら、評価制度では「自分でやってください」と伝えていることになります。人は、言われたことだけではなく、実際に評価されることへ適応します。

だから、**Role と KPI の整合を取る必要があります。**

大きな目標を、次の一週間まで下ろす

ここまでで、佐藤さんの最初の Plan の骨格ができました。

本人が目指しているのは、海外拠点責任者になることです。その意味を掘り下げると、現地社員を育て、自分がいなくてもチームとして成果を出せる経営人材になることでした。

会社は海外事業を拡大したい。現在のチームは、現地社員だけで一定範囲の顧客対応を完結できる状態を目指している。そこで佐藤さんには、現地社員へ判断を移し、自律を支援する Role が設定され、その成果を確認するための KPI も置かれます。

しかし、ここまでではまだ足りません。明日から、佐藤さんは何を変えるのでしょうか。

「部下を育てる」「もっと仕事を任せる」だけでは、まだ抽象的です。

そこでAIが、「来週、普段なら自分で答えを出している場面の中で、一つだけ違う行動を試すとしたら、何ができますか。」と問いかけます。

佐藤さんは考え、次のように決めます。

「一つの顧客案件について、自分で答えを出す前に、まず現地社員へ本人の案を聞いてみます。」

会社の「海外事業を拡大する」という大きな戦略が、本人の将来像へつながり、チーム目標へつながり、Roleへつながり、最後には、**「今週、一度だけ、自分が答えを言う前に、本人の案を聞く」**という一つの行動まで下りてきました。

これが、佐藤さんの最初の週次行動(Weekly Action)です。

図 8-1

「何になりたいか」から Weekly Action まで

— 本人の未来と会社の未来を、一週間の行動へつなぐ —



Weekly Action は、評価のためではなく、Role を試すためにある

第6章で見たように、週次行動(Weekly Action)は人事評価のための項目ではありません。Role を実際の仕事の中で試すための、小さな学習実験です。

もし「Weekly Action を 100% 実行した人ほど高く評価する」という運用にすれば、本人は失敗しそうな難しい行動ではなく、確実に達成できる行動を選ぶようになります。それでは、学習のための実験ではなく、評価を下げないための行動になります。

佐藤さんが今回の行動を実行できなかったとしても、それだけで「育成意欲が低い」と評価するわけではありません。

なぜできなかったのか。どの場面で元の行動へ戻ったのか。何が本人を止めたのか。そこにも、次の学習につながる情報があります。

成果(Result)には責任を持つ。しかし、そこへ向かう過程では、小さく試し、失敗からも学ぶ。この二つを分けることが重要です。

最初の Plan は、正解ではなく仮説である

ここまで丁寧に Plan をつくっても、それが正しいとは限りません。

佐藤さん自身が、本当に海外拠点責任者を目指したいのか。実際に管理職の仕事を経験する中で、「自分は専門家として生きたい」と気づく可能性もあります。

上司が設定した Role が適切ではないこともあります。設定した KPI が、本当にその Role の成果を表していないこともあります。会社の戦略そのものが変わることもあるでしょう。

最初の Plan は、本人を縛る契約ではありません。**学習を始めるための仮説です。**

仮説だから、仕事の中で試します。試した結果を見ます。違っていれば修正します。本人が成長すれば、ありがたい姿そのものが変わることもあります。

Plan もまた、学習の対象なのです。

AI がすること、人間がすること

ここまでの流れを見ると、AI と人間の役割の違いも明確になります。

AI は、本人へ問いを返します。「海外拠点責任者になりたい」という曖昧な言葉を掘り下げ、本人の未来と会社の方向との接点を考えるための材料を整理します。Role や KPI を考えるときには、見落としている可能性のある視点を示すこともできます。

しかし、本人が何をを目指すのかを決めるのは本人です。会社がどこへ向かうのかを決めるのは会社です。チームの中で誰ほどの Role を任せるのかを判断するのは上司です。そして、次にどの行動を試すのかは、本人自身の選択を含めて決めます。

AI に人を判断させるのではありません。AI が問い、整理し、選択肢を広げる。人間が意味づけし、選び、判断し、責任を持つ。

これが、EvaliA における人間参加型(Human in the Loop)の Plan 段階での具体的な姿です。

Plan は、一週間後に初めて意味を持つ

佐藤さんには、最初の週次行動(Weekly Action)が決まりました。

「今週、一つの顧客案件について、自分で答えを出す前に、まず現地社員へ本人の案を聞いてみる。」

しかし、これを決めただけでは、まだ何も変わっていません。一週間後、佐藤さんは本当に実行したのでしょうか。実行したなら、何が起きたのでしょうか。できなかったなら、どこで止まったのでしょうか。

そして、その経験によって、「人に任せる」という佐藤さん自身の見方は変わったのでしょうか。

人は、目標シートを書いた瞬間に成長するわけではありません。実際の仕事の中で何かを試し、経験し、振り返り、次の行動を変えることで学びます。

Plan は、実行されて初めて学習になります。ここから、EvaliA の日常運転が始まります。

次章では、佐藤さんの最初の一週間を追いながら、週次人工知能(AI)コーチングが、一つの小さな行動をどのように学習へ変えていくのかを見ていきます。

第9章 毎週、AIは何を問いかけるのか

——週次 AI コーチングと行動過程(Process)

第8章で、佐藤健太さんの最初の計画(Plan)ができました。将来は海外拠点責任者になりたい。その意味を掘り下げると、「現地社員を育て、自分がいなくてもチームとして成果を出せる経営人材になりたい」という姿が見えてきました。

会社が目指す海外事業の拡大とも接点があり、上司からは、「現地社員へ判断を移し、本人たちだけでも一定範囲の仕事を完結できる状態をつくる」という役割(Role)が置かれました。

そして佐藤さんは、最初の週次行動(Weekly Action)として、「今週、一つの顧客案件について、自分で答えを出す前に、まず現地社員へ本人の案を聞いてみる」と決めました。

しかし、目標を決め、Roleを決め、KPIを設定しただけでは、人はまだ変わっていません。重要なのは、そのPlanが実際の仕事の中でどのような行動になり、そこで何を経験したかです。

一週間後、EvaliAの日常運転が始まります。

一週間後、前回の Weekly Action へ戻る

AIとの週次コーチングは、毎回まったく新しい質問をするところから始まるわけではありません。まず、前の週に本人が何を試すと決めたのかへ戻ります。

AIは佐藤さんへ、「前回、『一つの顧客案件について、自分で答えを出す前に、まず現地社員へ本人の案を聞いてみる』と決めました。実際に試すことはできましたか」と問いかけます。

佐藤さんは、「はい。一度やってみました」と答えました。ここで「達成」と記録して終われば、これはタスク管理です。

しかし、EvaliAが見たいのは、約束を守ったかどうかだけではなく、その行動によって、実際の仕事の中で何が起きたのかです。

そこでAIは、「どのような場面でしたか」「相手はどのように答えましたか」「そのとき佐藤さん自身は、どのように反応しましたか」と、経験の中身を掘り下げていきます。

週次コーチングの中心にあるのは、できたかどうかではなく、**やってみた結果として、本人が何を経験したのか**という問いです。

「答えを教えない」という小さな実験

佐藤さんは、ある顧客から現地社員へ質問が来たときのことを話しました。

これまでなら、その質問を見た佐藤さん自身が回答を考え、現地社員へ「こう答えてください」と伝えていました。しかし今回は、Weekly Actionを思い出し、「あなたなら、どう回答しますか」と先に聞いてみました。

現地社員は少し考え、自分なりの回答案を出しました。佐藤さんから見ると、その答えには足りない点があります。いつもの佐藤さんなら、すぐに修正した答えを教えていたでしょう。

しかし今回は、もう一度、「なぜそう考えたのですか」と聞いてみました。すると現地社員は、自分の考えを説明している途中で、最初の回答案に足りなかった点へ自分で気づきました。

AIは、その場面を振り返りながら、「もし最初に佐藤さんが正しい答えを教えていたら、何が違っていたと思いますか」と問いかけます。

佐藤さんは、「その場では早く終わったと思います。でも、本人が自分で考える機会はなかったと思います」と答えました。

これは非常に小さな出来事です。しかし、佐藤さんの中では重要な経験が生まれています。

これまで佐藤さんは、「正しい答えを早く教えることが、部下を助けることだ」と考えていました。しかし今回、「答えを教えない方が、本人が考えられる場合もある」と実際に経験しました。

AIが佐藤さんへ「正しい管理職像」を教えたわけではありません。本人が実際の仕事で行動を変え、その結果を問いによって振り返ったことで、自分自身の見方を問い直す材料が生まれたのです。

うまくいったときにも、「なぜ」を聞く

人材育成では、できなかったことばかりに目が向きがちです。しかし、うまくいったときにも重要な学習があります。

そこでAIは、「今回うまくいった理由は何だと思えますか」「別の社員でも同じ方法を使えそうですか」「どの仕事なら使えて、どの仕事では難しいと思えますか」と問いを続けます。

佐藤さんは、「今までは、任せるか、自分でやるかの二つで考えていた気がします。全部任せなくても、まず本人に考えてもらうことはできます」と答えました。

ここで、佐藤さんの中の「任せる」という言葉の意味が変わり始めます。

以前は、「全部任せる」か「自分でやる」という二択に近かったものが、「本人に考えてもらう」「理由を聞く」「途中まで任せる」「最後だけ自分が確認する」といった複数のやり方に分かれてきます。

これは、単に Weekly Action を一つ達成したという話ではありません。仕事を見る枠組みそのものが少し広がっています。第4章で扱ったメンタルモデル(Mental Models)への気づきが、実際の仕事の中で始まっているのです。

次の週、佐藤さんは実行できなかった

ところが、その次の週はうまくいきませんでした。重要な顧客から難しい質問が届きます。佐藤さんは、現地社員へ一度考えてもらおうとしましたが、

「この顧客で失敗するわけにはいかない」と考え、結局、自分で回答をつくって顧客へ送りました。

一週間後の週次コーチングで、AIがWeekly Actionについて聞くと、佐藤さんは「今回はできませんでした」と答えます。

ここでEvaliAが「未達」とだけ記録し、実行率を下げるのであれば、重要な情報を失います。

AIは、「実行しようと思った場面はありましたか」「そのとき、何が止めましたか」と聞きます。

佐藤さんは、「重要な顧客だったので、失敗させるわけにはいかないと思いました。結局、自分で全部答えました」と話しました。

この回答から見えてくるのは、単なる行動不足ではありません。佐藤さんには、「**重要な仕事ほど、自分でやらなければならない**」という前提がある可能性があります。

そこでAIは、「重要な案件を全面的に任せることと、本人へ最初の案を考えてもらうことは同じでしょうか」と問い返します。

佐藤さんは少し考え、「確かに違います。最終判断は自分がするとしても、本人に考えてもらうことはできました」と答えます。

Weekly Action そのものは実行できませんでした。しかし、学習としては必ずしも後退ではありません。むしろ、**なぜ自分が元の行動へ戻ったのかが見えたことで、次の学習材料が生まれたのです。**

行動過程(Process)は、達成率ではない

本書でいう行動過程(Process)とは、Weekly Actionを何%達成したかという数字ではありません。

本人が何を試そうとしたのか。実際には何をしたのか。どこで前進し、どこで止まったのか。そこにはどのような阻害要因があったのか。本人は経験から

何に気づいたのか。そして、次の一週間に何を試そうとしているのか。その一連の変化を見ます。

同じ「未達」でも、意味はまったく違います。単に忘れていたのかもしれませんが。時間がなかったのかもしれませんが。失敗が怖かったのかもしれませんが。上司から止められたのかもしれませんが。そもそも本人には変えられない制度が障害になっていた可能性もあります。

それらをすべて「未達=0」としてしまえば、学習に必要な情報が消えてしまいます。

だから、Weekly Action そのものを人事評価の点数にはしません。もし実行率が高い人ほど高く評価されるなら、本人は難しいことへ挑戦せず、確実に達成できる行動を選ぶようになるからです。

Weekly Action は評価項目ではなく、本人が成長のために試す小さな実験です。

AI が「時間軸を持った鏡」になる

一度の対話であれば、人間の上司やコーチでもできます。AI の意味がより大きくなるのは、対話が毎週続き、時間軸ができたときです。

例えば三か月前、佐藤さんが「部下にはまだ判断できないので、自分が答えを出した方が早いです」と話していたとします。しかし最近は、「まず本人の案を聞き、必要なところだけ自分が修正しています」という発言が増えている。

AI は、その変化を本人へ返すことができます。

「三か月前には、『自分が答えを出した方が早い』と話していました。最近では、『本人の案を先に聞く』という行動が増えています。この変化を、自分ではどう感じていますか。」

本人は日々の仕事の中では、自分の変化に気づいていないかもしれませんが。しかし、過去と現在を並べることで、「確かに以前より任せられるようになっている」と認識できます。

一方、逆のケースもあります。佐藤さん自身は、「かなり任せられるようになった」と感じている。しかし数週間の対話を見ると、通常案件では任せるようになったものの、重要案件では毎回自分が最終回答まで行っている。

その場合には、「通常案件では本人へ任せる行動が増えています。一方、重要案件では五週続けて佐藤さん自身が最終回答まで行っています。この違いをどう考えますか」と問い返すことができます。

本書では、こうして過去の自分と現在の自分を比較し、本人だけでは気づきにくい変化や繰り返しを鏡として返すことを、「**内省支援フィードバック**」(**Reflective Feedback**)と呼びます。

AIが「あなたはこういう人だ」と決めない

時間軸の情報が増えると、AIはさまざまなパターンを見つけられるようになります。しかし、ここで注意しなければならないことがあります。

例えば、重要案件になると佐藤さん自身が判断しているという事実から、AIが「あなたは部下を信用していません」と断定することです。

しかし、理由が本当に「部下を信用していないから」とは限りません。契約上、佐藤さん自身に最終責任があるのかもしれませんが。顧客との約束上、判断権限を渡せない案件なのかもしれません。

そこでAIは、「通常案件と重要案件で行動が違います。重要案件では何が違うのでしょうか」と問い返します。

本人が、「やはり自分が失敗を怖がっているのだと思います」と考えることもできますし、「この案件では権限上、自分が判断する必要があります」と訂正することもできます。

AIの役割は、人間の内面を判定することではありません。**見えている事実から仮説をつくり、本人が考えるための問いとして返すこと**です。

本人だけを見ていると、原因を間違える

週次コーチングを続けていると、佐藤さんが「現地社員がなかなか主体的に動いてくれません」と話すことがあるかもしれません。

佐藤さん本人だけを見ていれば、「現地社員の主体性を高めるにはどうすればよいか」という話になりそうです。しかし AI は、別の視点から問いを返すことができます。

現地社員から見ると、どこまで本人が判断してよいか明確なのか。本人が判断して失敗した場合、どのような扱いになるのか。佐藤さんがいない場面でも、本当に本人は判断できていないのか。他の現地社員にも同じことが起きているとすれば、個人以外の原因は考えられないのか。

こうした問いによって、「現地社員が主体的ではない」という個人の問題が、「権限が曖昧である」「上司が判断を取りすぎている」「失敗した場合の責任範囲が不明確である」といった構造上の問題へ変わることがあります。

AI が原因を決めるのではありません。本人、上司、チーム、制度など、**原因を探す視野を広げる**のです。

最後に、次の一週間で何を試すかを決める

振り返りがどれほど深くても、次の行動につながらなければ学習循環は閉じません。そこで毎週の対話の最後には、「では、次の一週間で何を試しますか」という問いへ戻ります。

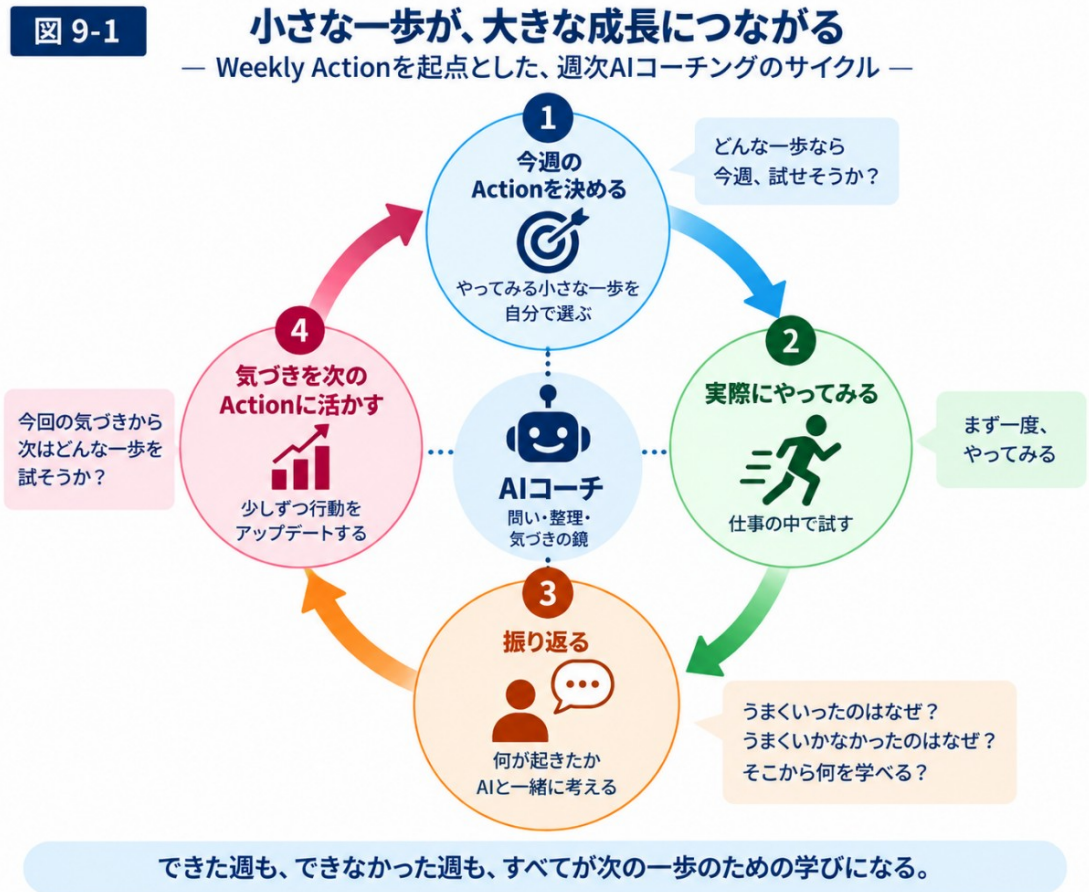
今回、佐藤さんは、「通常案件では任せられるようになったが、重要案件になると自分で全部やってしまう」というパターンに気づきました。

そこで次の Weekly Action として、「重要案件でも、最終判断をする前に、一度は現地社員へ本人の案と理由を聞く。」と決めたとします。

これは「もっと部下を信頼する」という精神目標ではありません。翌週、実際に試したかどうかを振り返ることのできる具体的な行動です。

AI が候補を示すことはできますが、最終的に何を試すのかは本人が選びます。もちろん、法令、安全、セキュリティ、会社として守るべき手順は自由選択の対象ではありません。

守るべきものは会社が決める。一方、成長のために試す行動には本人の選択を入れる。この区別は変わりません。



一週間の対話から、何が残るのか

週次コーチングが終わると、単なる会話ログだけが残るわけではありません。

前回の Weekly Action、実際に起きたこと、前進や停滞、主な阻害要因、本人が気づいたこと、そして次の Weekly Action といった学習情報が構造化されます。

これらが毎週蓄積されることで、一回の対話では分からなかったパターンが見えるようになります。

例えば、「三週間、同じところで止まっている」「以前は上司へすぐ答えを求めていたが、最近は自分の案を持って相談するようになった」「通常業務では行動が変わったが、重要案件だけ元の行動へ戻っている」といった変化です。

こうした情報は、本人を管理するための履歴ではありません。本人自身が、自分の変化や繰り返しに気づくための学習材料です。

AIが大きく変えるのは、一度だけ優れた答えを返せることではありません。問いと振り返りを、仕事の中で継続できることです。

本人では変えられない問題は、本人へ返し続けない

佐藤さんが何週間か Weekly Action を続けていると、別の問題が見つかることがあります。

本人は現地社員へ判断を任せたい。しかし会社の決裁規程では、小さな案件であっても日本人責任者の承認が必要になっている。この状態で佐藤さんへ毎週、「もっと任せましょう」と問い続けても意味がありません。問題は、本人の成長ではなく仕組みにあります。

そこで、本人では変えられない構造上の問題が見えた場合には、週次コーチングから SCAD へつなぎます。

本人で変えられることは、次の Weekly Action へ戻す。一方、本人だけでは変えられないことは、チームや組織の学習へ上げる。

例えば、「重要案件でも先に本人の案を聞く」は佐藤さん自身で試せます。しかし、「会社の決裁権限を変更する」は佐藤さん一人ではできません。

EvaliA が個人の気づきを生み、SCAD が必要な気づきを組織改善へつなげる。この分岐が、個人学習を組織学習へ広げるうえで重要になります。

対話のすべてを上司へ渡すわけではない

この週次対話が成立するためには、もう一つ重要な条件があります。

例えば佐藤さんが AI へ、「実は、上司の方針には少し疑問があります」と話したとします。

その文章がそのまま上司へ送られると分かっているならば、本人は次から同じように話しにくくなるでしょう。AI との対話が、自分を振り返る場ではなく、上司への報告書になってしまうからです。

よって、本人が自分を振り返るための生の対話と、組織が学習のために使う情報には境界が必要です。組織に必要なのは、本人が話したすべての文章ではありません。「同じ阻害要因が複数週続いている」「Role と権限にずれがある可能性がある」「構造上の問題が提起されている」といった、組織が学習するために必要な情報へ整理して扱うことができます。

何を本人の内省空間に残し、何を本人が共有し、何を組織学習へ使うのか。そのルールを明確にしておくことが、観測を監視へ変えないために必要です。

小さな一週間を繰り返す

佐藤さんは、一週間で別人になるわけではありません。

最初の週には、一度だけ現地社員へ先に考えてもらった。次の週には、重要案件で元の行動へ戻った。その経験から、「重要な仕事ほど自分がやらなければならない」という自分の前提に気づいた。そして次の週には、最終判断は自分が持ちながらも、先に本人の案を聞いてみる。

また別の週には、うまくできないこともあるでしょう。それでも振り返り、次の一週間でもう一度試します。この小さな循環が続くことで、数か月後には、「以前より、自分が答えを出さなくても仕事が進む場面が増えた」という変化が生まれるかもしれません。

人が成長するのは、一度の優れた研修を受けた瞬間ではありません。仕事の中で試し、経験し、考え直し、もう一度試すことを続けた結果です。

AI の役割は、その学習を一度だけ起こすことではありません。**学習を、毎週途切れさせないこと**にあります。

一人が学ぶだけでは、まだ学習する組織ではない

ここまでで、佐藤さん個人には学習循環が生まれました。

しかし、佐藤さんだけが成長しても、それだけでは学習する組織(Learning Organization)にはなりません。同じチームには、佐藤さん以外の社員もいます。それぞれが Weekly Action を持ち、仕事を振り返っていると、複数人の間に共通するパターンが見えてくることがあります。

一人だけが「判断権限が分からない」と話しているのか。それとも、何人も同じところで止まっているのか。成果へ向かう行動は多いが、安全確認が不足していないか。新しい提案は出ているのに、実行へ移す行動が弱くないか。

一人の対話だけでは見えなかったものが、チームを見ることで初めて見えてきます。ここから、学習の単位を一段上げます。

次章では、佐藤さんを含む一人ひとりの週次学習を、本人の生の対話をそのまま公開することなく、どのように**チーム学習(Team Learning)**へ変えていくのかを見ていきます。

そこで中心となるのが、チーム CAGE マップ(Team CAGE Map)とチームフィードバック(Team Feedback)です。

第10章 個人の学びを、どうチームの学びへ変えるのか

——チーム CAGE マップ(Team CAGE Map)とチームフィードバック(Team Feedback)

第9章では、佐藤健太さんが毎週の AI コーチングを通じて、自分の行動を振り返る様子を見てきました。「部下へ仕事を任せる」という抽象的な目標を週次行動(Weekly Action)へ変え、実際の仕事の中で試す。うまくいったことも、できなかったことも行動過程(Process)として振り返る。その中で佐藤さんは、「重要な仕事ほど自分で判断しなければならない」という自分自身のメンタルモデル(Mental Models)にも少しずつ気づき始めました。

この循環を続ければ、佐藤さん自身は成長していくでしょう。しかし、佐藤さん一人が成長しても、それだけでは学習する組織(Learning Organization)にはなりません。同じチームには複数の社員がいて、それぞれが AI との対話を続け、それぞれの仕事を振り返っています。

一人ひとりが学んでいるだけなら、それは「複数の個人学習」が並んでいる状態です。ここから必要になるのは、**個人の学びの中から、チーム自身が学べる情報を取り出す**ことです。

佐藤さんだけを見ていると、「佐藤さんの問題」に見える

ある週、佐藤さんが週次 AI コーチングの中で、「現地社員へもっと判断を任せたいのですが、結局こちらへ確認が来てしまいます」と話したとします。

佐藤さん一人だけを見れば、「佐藤さんの任せ方に問題があるのではないか」と考えるかもしれません。実際、第9章では佐藤さん自身にも、「重要な

仕事ほど自分で判断しなければならない」という前提がある可能性が見えてきました。

ところが、同じチームの別の日本人社員も、「現地社員から細かい確認が多く、こちらの判断待ちで仕事が止まります」と話していた。さらに別の社員も、「自分で考えてもらいたいのですが、最後は日本人側への確認になります」と振り返っていたとします。

一人だけなら、個人の問題かもしれません。しかし、複数人が同じところで止まっているなら、見方を広げる必要があります。

現地社員の能力が足りないのかもしれませんが、しかし、判断権限そのものが曖昧なのかもしれない。失敗した場合に誰が責任を持つのか分からないのかもしれない。「最後は日本人が決める」という暗黙の習慣があるのかもしれませんが。会社の決裁ルール自体が、現地社員へ判断を移す設計になっていない可能性もあります。

個人だけを見ていると「その人の問題」に見えていたものが、複数人を並べた瞬間に、チームや仕組みの問題として見えてくることがあります。

これが、個人学習をチーム学習(Team Learning)へ変える最初の一步です。

だからといって、全員の対話を上司へ見せるわけではない

では、チーム全体の状態を知るために、一人ひとりがAIと話した内容を上司が全部読めばよいのでしょうか。それでは、学習のための対話そのものが成り立たなくなる可能性があります。

例えば佐藤さんがAIとの対話の中で、「上司が重要な案件では最後に全部判断してしまうので、本当に自分へ権限があるのか分かりません」と話していたとします。別の社員は、「会議で意見を言っても、最後には上司の案になります」と話しているかもしれません。

こうした生の対話が、そのまま上司へ伝わると分かっているならば、社員は次第に本当に考えていることを話さなくなるでしょう。本人が自分を振り返るための場が、上司へ提出する報告書へ変わってしまうからです。

そこで、本人の生の対話と、チームが学習するための情報を分けます。AIとの対話そのものを公開するのではなく、そこから「判断権限が不明確という阻害要因が複数人に見られる」「現地社員へ任せようとする行動はあるが、最終的に日本人側へ判断が戻るケースが繰り返されている」「新しい意見を出す行動より、既存方針に従う行動が多い」といった、チームが学ぶための行動パターンへ構造化します。

つまり、**本人の生の対話から、チームが学べる行動パターンへ変換する**のです。

個人の対話から、チームの情報へ

情報の流れを整理すると、最初にあるのは本人とAIとの対話です。そこから、本人が何を試したのか、何が前進したのか、どこで止まったのか、どのような阻害要因があったのか、Roleとの間にどのようなずれがあるのか、といった個人レベルの学習情報を構造化します。

次に、複数人の情報を比較します。一人だけに出ていることなのか。複数人に共通しているのか。同じ種類の行動が繰り返されているのか。チームとして一定の偏りが生まれているのか。

そして、その結果をチームフィードバック(Team Feedback)へ変えます。

全体の流れは、**本人の生の対話 → 個人レベルの学習情報 → 複数人の共通パターン → Team Feedback**となります。

図10-1 個人学習からチーム学習へ

個人の対話を、チームの学習と行動に変える流れ



ここで重要なのは、AIが扱える情報と、上司が見るべき情報は同じではないということです。

特に人数の少ないチームでは、集約したつもりでも「これは誰の発言か」が推測できてしまうことがあります。そのため実際の運用では、どの人数から集約表示するのか、誰が閲覧できるのか、本人が再特定されるような情報をどこまで出すのか、といったルールを明確にする必要があります。

学習するために情報を使うことと、人を監視することは違います。

Team CAGE Map で見るのは、「誰が何タイプか」ではない

複数人の実際の行動を見るとき、本書ではCAGE適応モデル(CAGE Adaptation Model)を使います。

第5章で説明したとおり、見るのは**成果(Challenger)**、**探究(Adventurer)**、**安全(Guardian)**、**関係(Empathizer)**という四つの行動機能です。

ただし、ここで使うチーム CAGE マップ(Team CAGE Map)は、「佐藤さんは Challenger 型」「田中さんは Guardian 型」と人を分類するためのものではありません。また、自己診断結果を何人分か平均して、「このチームは成果 40%、安全 20%」と決めるものでもありません。

見るのは、**実際の仕事の中で、どの行動機能が使われているか**です。

人は状況 (Context) によって行動を変えます。普段は慎重な人でも、新しい顧客を前にすれば探究(Adventurer)の行動を多く使うことがあります。人との関係を大切にしている人でも、納期が迫れば成果(Challenger)の行動を強く使うかもしれません。

したがって、本人が自分をどう認識しているかと、実際の仕事でどの機能を使っているかは、別の情報として見ます。

佐藤さんのチームには、どんな行動が起きているのか

佐藤さんのチームには、日本人社員と現地社員を含む数名のメンバーがいます。チームの目標は、「現地社員だけで一定範囲の顧客対応を完結できる状態をつくる」ことでした。

数週間分の週次学習からチーム全体の行動を見ると、まず成果(Challenger)に関する行動は多く見られました。顧客案件を前へ進める。期限を決める。数字を追う。問題が起きればすぐ対応する。チーム全体として、「結果を出す」という方向には強く動いています。

探究(Adventurer)についても、新しい顧客への提案方法を試したり、現地市場に合ったサービスを考えたりする行動が見られます。

一方、安全(Guardian)では、新しい案件を急ぐあまり、契約条件や回収リスクの確認が後回しになるケースがありました。また、関係(Empathizer)を見る

と、日本人メンバー同士では頻繁に相談している一方、現地社員自身から意見を引き出す行動は少ない。会議には参加していても、実際の判断は日本人側で行われることが多いというパターンが見えてきました。

ここで、「佐藤さんは関係(Empathizer)が弱い」とは判断しません。問うべきなのは、**現在のチームでは、どの行動が起きやすく、どの行動が起きにくいのか**です。

Team CAGE Map は、成績表ではなく「チームの鏡」である

Team CAGE Map を見ると、「成果へ向かう行動が多い」「安全確認の行動が少ない」「現地社員から意見を引き出す行動が少ない」といった特徴が見えてきます。

しかし、それをチームの性格として固定してはいけません。事業の局面によって、必要な機能は変わるからです。

新規事業の立ち上げ時期なら、成果(Challenger)や探究(Adventurer)が多くなることは自然です。品質問題が続いているなら、安全(Guardian)を意識的に強める必要があります。国籍や部署を越えた協働が重要になっているなら、関係(Empathizer)の重要性が高まります。

四つの機能を同じ割合で使うことが目的なのではありません。重要なのは、**「現在のチーム目標と状況（Context）に対して、必要な機能が使えているか」**です。

その意味で、Team CAGE Map はチームの評価表ではありません。個人に人工知能(AI)が鏡を返したのと同じように、チーム全体へ「現在、自分たちはどのように働いているのか」という鏡を返すために使います。

佐藤さん個人の問題だと思っていたことが、チームの問題に変わる

第9章では、佐藤さん自身に「重要な仕事ほど自分でやらなければならない」というメンタルモデル(Mental Models)がある可能性が見えました。

しかし、チーム全体を見ると、佐藤さんだけではありませんでした。日本人メンバーの多くが、重要な判断になるほど自分で答えを出しています。現地社員から意見を聞く行動も少ない。そして会社の承認ルール自体も、日本人責任者による最終確認を前提としている部分がある。

ここまで来ると、「佐藤さんがもっと部下を信頼すればよい」だけでは解決しません。個人のメンタルモデルに加えて、会議の進め方、Role、権限、会社の決裁ルールといった構造が、同じ行動を生み出している可能性があります。

個人を深く見ることと、構造を見ることは対立しません。むしろ、**個人の学習を複数人並べたときに、初めて構造が見えてくる**のです。

上司は「誰を直すか」ではなく、「チームに何を返すか」を考える

Team CAGE Map や共通する行動パターンを見た上司は、どうすればよいのでしょうか。

ここで、「佐藤さんは、もっと現地社員との関係を大切にしてください」と個人へ戻してしまえば、せっかくチームまで広げた学習が、また一人の人材教育へ縮小してしまいます。

例えば上司は、チーム全体へこう返すことができます。

「最近の行動を見ると、顧客案件を前へ進める行動は非常に多い一方で、現地社員から意見を引き出したり、本人に判断してもらったりする行動は少ないようです。これはなぜだと思いますか。」

これがチームフィードバック(Team Feedback)です。

「このチームは関係(Empathizer)が弱い」と断定するのではなく、「**こういう行動の偏りが見えている。自分たちはどう考えるか**」とチーム自身へ問いを返します。

AIはパターンを整理することができます。上司は問いを返すことができます。しかし、そのパターンにどのような意味があるのかを考えるのはチーム自身です。

チームの中に、違う意見があってもいい

Team Feedbackを受けたからといって、全員が同じ結論になる必要はありません。

佐藤さんは、「もっと現地社員へ任せられると思います」と考えるかもしれません。別の社員は、「まだ経験が不足しているので、これ以上任せるのは危険です」と反対するかもしれません。さらに別の社員は、「能力の問題ではなく、どこまで判断してよいか本人たちが分かっていないのではないか」と考えるかもしれません。

この違いには意味があります。成果(Challenger)の視点からは、もっと早く前へ進めたい。探究(Adventurer)の視点からは、小さな範囲で新しい任せ方を試してみたい。安全(Guardian)の視点からは、失敗のリスクと責任範囲を確認したい。そして関係(Empathizer)の視点からは、現地社員本人がどう感じているのかを聞く必要がある。

どれか一つを正解にするのではなく、異なる視点を同じ問題へ当てます。違いがあることが問題なのではありません。違いをチームの判断能力として使えないことが問題なのです。

良いチームは、対立しないチームではない

例えば佐藤さんが、「まず小さく任せてみましょう」と提案したとします。

別の社員が、「その前に契約上のリスクを確認した方がよい」と反対しました。

これを人格の違いとして見れば、「佐藤さんは強引だ」「相手は慎重すぎる」という対立になります。

しかし、機能として見れば違います。佐藤さんは成果(Challenger)や探究(Adventurer)を使って前へ進めようとしている。一方の社員は、安全(Guardian)の機能を使って見落としを防ごうとしている。

すると議論は、「どちらが正しいか」から、「今回の Context では、前へ進めることとリスク確認を、どの順番で組み合わせるか」へ変わります。

チーム学習(Team Learning)とは、対立をなくすことではありません。異なる機能を、より良い判断のために使えるようにすることです。

Team Feedback を、Team Action へ変える

チームの状態について話し合っただけでは、実際の仕事は変わりません。

そこで、個人の Weekly Action と同じように、チームにも小さな実験を置きます。

佐藤さんのチームでは、「現地社員から意見を引き出す行動が少ない」というパターンが見えていました。ここで、「もっと現地社員を尊重しよう」と決めても、翌日の行動にはなりません。

代わりに、「次回の定例会議から、日本人メンバーが意見を言う前に、現地社員から先に意見を聞く。」という具体的な Team Action を決めます。

一週間後、現地社員の発言は増えたのか。会議の内容は変わったのか。判断は遅くなったのか。それとも、これまで日本人側に見えていなかった情報が出るようになったのか。実行した結果を、もう一度観測します。

チームでも、Team Feedback→Team Action→実行→再観測という短い学習循環を回すのです。

行動を変えるだけでなく、Role を変えることもある

チームの問題によっては、Team Action だけでは足りません。

例えば、安全(Guardian)に関する行動が不足していることが分かったとします。全員へ「もっとリスクを意識しましょう」と伝えるだけでは、誰が責任を持つのか曖昧なままです。

そこで、新規契約については特定の社員がリスク確認を主に担うなど、Role そのものを調整することがあります。

ただし、その社員だけが安全について考えればよいという意味ではありません。誰が主に担うのかを明確にしながら、チーム全体として必要な機能を失わないようにします。

また、状況 (Context) が変われば Role も変えます。

Team CAGE Map は、人を一つの役割へ固定するためのものではありません。**現在のチーム目標と Context に合わせて、Role と協働関係を調整するための材料として使います。**

複数人が同じところで止まれば、共通 Gap として見る

チームを見ると、CAGE とは別に重要な情報が見えてきます。それは、複数人が同じ場所で止まっているという情報です。

例えば佐藤さんを含む複数の日本人社員が、「現地社員へ判断を任せようとしたが、最終的には日本人責任者の承認が必要だった」と話しているとします。

一人だけなら、本人の仕事の進め方かもしれません。しかし、複数人が同じところで止まっているのであれば、共通 Gap として見る価値があります。

会社の決裁ルールが現地への権限移譲と矛盾しているのかもしれません。Role の定義が曖昧なのかもしれません。現地社員へ責任だけを与え、権限を十分に渡していない可能性もあります。

このような問題は、チームの中だけでは変えられないことがあります。その場合には SCAD へつなぎます。

本人で変えられるなら Weekly Action。チームで変えられるなら Team Action や Role 調整。そして、組織の仕組みを変える必要があるなら SCAD へ上げる。

問題の場所によって、学習の出口を変えるのです。

一人の成功も、チーム学習へ変えられる

チーム学習へ変えるのは、問題だけではありません。成功も同じです。

第9章で佐藤さんは、「答えを教える前に、本人の案と理由を聞く」という方法が有効であることに気づきました。佐藤さんだけが使っていれば、個人の学習です。

しかし Team Feedback の場で、「他の管理職も一度試してみよう」となれば、個人の経験がチームの実験へ変わります。数週間後、他のメンバーでも、現地社員が自分の案を持って相談する場面が増えたとします。

ここで初めて、一人の成功がチームの能力へ変わり始めます。

ただし、一度うまくいったからといって、すぐに全社標準にはしません。本当に他の人にも再現できるのか。どの Context なら有効なのか。どこまで標準化できるのか。

この問題は、第14章で扱う名人循環(Mastery Loop)へつながっていきます。

上司自身も、チームの状態をつくっている

Team CAGE Map を見ていると、部下だけではなく、上司自身の行動が学習対象として見えてくることがあります。

例えば、チームで探究(Adventurer)の行動がほとんど出ていない。上司は、「うちの社員は新しい提案をしない」と感じている。

しかし、複数人の週次学習を構造化すると、「以前、新しい案を出したが、その場で上司から否定された」という経験が何度も出ているとします。

この場合、「社員の探究力が弱い」と結論づけるのは早いでしょう。上司自身の反応が、新しい提案を出しにくい環境をつくっている可能性があります。

人工知能(AI)は上司へ、「提案行動が少ない一方、提案後に否定された経験も複数見られます。上司自身の反応が、次の提案行動へどのような影響を与えている可能性がありますか」と問い返すことができます。

上司は、チームを外から観察している人ではありません。上司自身も、チームの状態をつくっている一人です。したがって、上司の行動もまた、チーム学習の対象になります。

Team CAGE Map を、人事評価にしない

ここには重要な注意があります。Team CAGE Map を見て、「安全(Guardian)が少ないから評価を下げよう」「関係(Empathizer)の行動が多い人を高く評価しよう」と使い始めたら、何が起きるでしょうか。

社員は、仕事のためではなく、評価のために「評価されそうな行動」を見せるようになります。

これは、Weekly Action を人事評価に直接使わないのと同じ理由です。

Team CAGE Map の目的は、誰を高く評価し、誰を低く評価するかを決めることではありません。**今のチーム目標に対して、自分たちはどの機能を使えていて、どこを変える必要があるのかを学ぶための鏡**です。

個人の成果評価に必要な情報は、後に成果(Result)として別の経路から確認します。学習のための情報と、評価のための情報を最初から一つにしない。この区別が重要です。

チーム学習(Team Learning)は、仲良くなることではない

チーム学習(Team Learning)という言葉から、メンバー同士が集まり、互いを理解し、仲良くなることを想像するかもしれません。もちろん、良い関係には価値があります。

しかし、本書で考えているチーム学習は、それだけではありません。

チームとして何を目標しているのか。実際にはどのような行動が起きているのか。必要な CAGE 機能が使われているのか。複数人が同じところで止まっていないか。Role や権限にずれはないか。上司自身の行動を含め、働き方のどこを変える必要があるのか。

そうしたことを、自分たち自身の学習対象にします。つまり、**チームが自分たちの働き方そのものを観測し、試し、変え、また観測することが、本書でいうチーム学習です。**

佐藤さんの課題は、もう佐藤さんだけの課題ではない

第8章の時点では、佐藤さんの課題は「自分で仕事をしすぎることに」見えていました。

第9章では、その背後に「重要な仕事ほど自分が判断しなければならない」という本人のメンタルモデル(Mental Models)が見えてきました。

しかし本章でチーム全体を見ると、さらに別のものが見えてきます。現地社員へ判断を任せにくいのは、佐藤さん一人の問題ではない。日本人メンバーの多くが同じ行動をしている。会議の進め方にも原因がある。Role や権限も曖昧である。そして、会社の決裁制度まで影響している可能性がある。

個人から始まった学習が、チームを通じて構造へ広がりました。ここで初めて、佐藤さんの成長が、佐藤さん一人のものではなくなります。

しかし、行動だけでは人をすべて理解できない

ここまで第9章、第10章では、主として行動を見てきました。佐藤さんが何を試したのか。どこで止まったのか。チーム全体ではどのような行動が起きているのか。CAGEの四つの機能が、実際の仕事の中でどのように使われているのかは重要です。

しかし、外から見える行動だけでは、本人がその仕事をどう受け止めているのかまでは分かりません。

佐藤さんは以前より現地社員へ仕事を任せられるようになっていきます。チームでも、現地社員自身が意見を話す場面が増えてきました。外から見れば、順調に成長しているように見えます。

ところが本人の中では、「自分が直接仕事をする量が減ってきた」「以前ほど自分自身が顧客から評価される場面がない」「部下が成果を出すことはうれしい。しかし、自分は会社からどう評価されるのだろうか」といった別の問いが生まれているかもしれません。

行動過程(Process)だけでは、この状態は分かりません。

そこで次章では、毎週の行動とは別に、一定の節目で、佐藤さん自身がこの期間をどう受け止めているのかを振り返ります。

AIが佐藤さんの心を診断するものではありません。本人自身が、何に達成感を感じ、何に不安や不満を感じ、何を成長と捉え、現在のRoleをどう受け止めているのかを言葉にします。

次章では、二つ目の観測経路である**心理状態(Mind)**を見ていきます。

第11章 節目では、本人は仕事をどう受け止めているのか

——心理状態(Mind)を点数ではなく対話で見る

第9章では、佐藤健太さんが毎週のAIコーチングを通じて、実際の仕事の中で何を試し、どこで止まり、何を学んだのかという行動過程(Process)を見てきました。第10章では、一人ひとりの行動をチームレベルで見ることによって、佐藤さん個人の問題に見えていたものの背後に、Roleや権限、会議の進め方、会社の仕組みといった構造が存在する可能性も見えてきました。

ここまで中心に見てきたのは、「何をしたのか」という行動です。

しかし、人は行動だけでは理解できません。同じように成果を出している二人でも、一人は「難しいけれど成長している」と感じ、もう一人は「この働き方を続けるのはつらい」と感じているかもしれません。反対に、まだ十分な成果が出ていなくても、本人自身は「今までで一番成長している」と感じている場合もあります。

そこでEvaliAでは、毎週の行動過程(Process)とは別に、四半期または半期など一定の節目で、本人自身がこの期間をどのように受け止めているのかを振り返ります。本書では、この観測経路を**心理状態(Mind)**と呼びます。

ただし、ここでいう心理状態(Mind)は、AIが本人の精神状態を診断するものではありません。医学的・臨床的な診断でもありません。見るのは、本人が仕事について語る満足、不満、達成感、負担、意味、成長実感などの自己認識です。

佐藤さんにとって、この半年はどういう時間だったのか

佐藤さんがEvaliAを使い始めてから、一定期間が経ちました。

最初は、「海外拠点責任者になりたい」という希望から始まりました。AIとの対話を通じて、その意味は「現地社員を育て、自分がいなくてもチームとして成果を出せる経営人材になりたい」という未来像へ深まりました。

その後、佐藤さんは毎週、小さな週次行動(Weekly Action)を試してきました。自分が答えを言う前に現地社員の案を聞く。すぐに修正せず、なぜそう考えたのかを聞く。小さな判断から任せてみる。重要案件では元のやり方へ戻ってしまうこともありましたが、そのたびに振り返り、次の行動を選び直してきました。

チームにも少しずつ変化が起きています。現地社員から自分の案が出る場面が増え、佐藤さんが直接答えを出さなくても仕事が進むケースも増えてきました。

行動過程(Process)だけを見るなら、佐藤さんは前進しています。

では、**佐藤さん自身は、この半年をどう受け止めているのでしょうか。**ここで、節目の振り返り対話が始まります。

「満足度は何点ですか」から始めない

社員の状態を知ろうとすると、会社は数字を取りたくなります。「仕事にどの程度満足していますか」「会社を友人へ勧めたいですか」といった質問にも意味があります。同じ形式で多くの社員から回答を集めれば、部署や時期による変化を比較しやすくなります。

例えば佐藤さんが満足度を「3」と答えたとして、その「3」だけでは、本人が何を感じているのかまでは分かりません。

従来型サーベイは、5段階、例えば、大変不満足、不満足、普通、満足、大変満足のように、最初から定量化した質問形式でした。

しかし、EvaliAを通じたサーベイは、あくまで「仕事の振り返り」として、社員はAIと定性的な対話を行うだけです。後に、AIが、管理者向けに社員の「動機づけ要因」、「衛生要因」、「エンゲージメント」のレベルを定量化させます。これがEvaliAの特徴の一つです。社員は、そもそも仕事の振り返り

をサーベイとも思わないので、一般的なサーベイよりも、管理者が実態や背景を把握しやすくなる可能性があります。

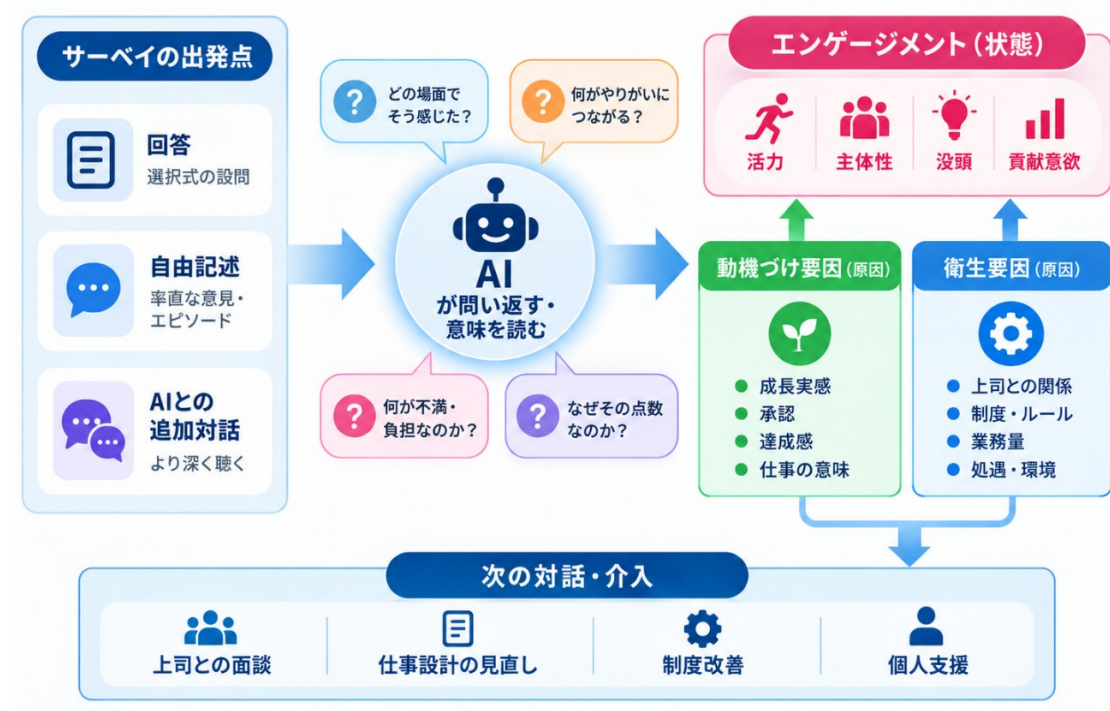
社員は、仕事内容には満足しているが、評価制度に不安があるのかもしれませんが。上司との関係は良いが、仕事が多すぎるのかもしれませんが。現在の仕事は苦しいけれど、自分自身は大きく成長していると感じている可能性もあります。

EvaliA では、定量化した点数だけで終わりません。本人がその状態を、どのような意味として受け止めているのか。そこを対話によって言葉にしていきます。

図 11-1

対話型サーベイ — 点数から意味へ

— AIとの対話から、エンゲージメントという状態と、その背景要因を読み解く —



「一番うれしかったことは何ですか」

AI は、佐藤さんへ例えばこう問いかけます。

「この期間を振り返って、一番うれしかったことは何ですか」

佐藤さんは、「現地社員が、自分で顧客への回答を考えてくれるようになったことです」と答えました。

以前なら、現地社員は「どう回答すればいいですか」と佐藤さんへ聞いていました。最近は、「私はこのように回答しようと思いますが、どうでしょうか」と、自分の案を持って相談する場面が増えています。

そこでAIは、「その変化の何が、佐藤さんにとってうれしかったのでしょうか」と問い返します。

佐藤さんは、「自分が全部やらなくても仕事が進むようになったことです。それに、現地社員自身が成長していると感じます」と答えます。

さらに、「それは、最初に話していた海外拠点責任者になりたいという未来と、どうつながっていますか」と聞かれ、佐藤さんは少し考えます。

「以前は、海外拠点責任者になるなら、自分が一番仕事のできる人でなければならなかったと思っていました。でも今は、自分がいなくても現地社員が動ける状態をつくる方が、責任者としては大事なのかもしれません。」

ここで行っているのは、満足度の測定ではありません。**本人が、自分の経験から「何を成長と考えるのか」を言葉にしているのです。**

ところが、佐藤さんには別の不安もあった

AIは、達成感だけでなく反対側からも問いかけます。

「では、この期間で一番不満だったこと、あるいは違和感があったことは何ですか。」

佐藤さんは少し考え、「自分自身が直接仕事をする量が減ったことです」と答えました。

現地社員が成長したことはうれしい。それなのに、自分が直接仕事をする量が減ったことには違和感がある。一見すると矛盾しているように見えます。

そこで AI は、「仕事量が減ったことそのものが不満なのでしょうか。それとも、別のことが気になっていませんか」と聞きます。

すると佐藤さんは、「自分の案件数や売上が減っています。部下が成果を出すことはうれしいのですが、自分自身が会社からどう評価されるのかが少し不安です」と答えました。

最初は「不満」と表現されていたものが、対話を続けると、実際には「**評価への不安**」であることが見えてきました。

この違いは重要です。

AI が最初の発言だけを見て「佐藤さんは会社に不満を持っている」と分類してしまえば、本人が実際に感じている意味を取り違える可能性があります。言葉を急いで分類するのではなく、本人自身に、その意味を確かめてもらいます。

達成感と不安は、同時に存在してよい

佐藤さんには、現地社員が成長していることへの達成感があります。同時に、自分が会社からどう評価されるのかという不安もあります。

この状態を、「満足している」「不満である」という一本の軸だけでは説明できません。

仕事そのものには意味を感じている。自分自身の成長も感じている。しかし、評価には不安がある。人の中には、こうした複数の状態が同時に存在します。

第 1 章で触れたハーズバーグの考え方も、この状態を考える一つの補助線になります。ただし、ここで重要なのは理論そのものではなく、**佐藤さんを一つの点数へまとめないことです。**

EvaliA では、何に達成感を感じているのか。何が不満なのか。何に負担を感じているのか。何を成長と感じているのか。それぞれを分けて見ます。

一つの数字にまとめてしまうと消えてしまう違いの中に、次の学習につながる情報があるからです。

なお、こうした定性的な対話から、「動機づけ要因」、「衛生要因」、「エンゲージメント」に関わる状態を、それぞれ複数の観測指標として構造化し、時系列で比較することはできます。

ただし、それは本人の心理状態を一つの総合スコアへ合算したり、心理尺度として精密に測定したりすることを意味しません。数値は、変化や違いを見つけるための補助情報として使い、その意味は本人の言葉と対話によって確かめます。

不満を、すぐに解消しようとしな

社員が不安や不満を語ると、会社はそれを「なくすべきもの」と考えがちです。

しかし、佐藤さんの評価への不安にも、少なくとも二つの可能性があります。

一つは、本当に会社側に矛盾がある場合です。会社は「部下を育ててください」と求めている。しかし実際の評価制度では、本人の売上だけを重視している。それなら、佐藤さんが不安になるのは自然です。

もう一つは、会社側ではチーム育成も評価しているにもかかわらず、佐藤さん自身が「自分で直接成果を出すことこそ、自分の価値だ」と考えている場合です。

前者なら、会社の制度や上司との認識を変える必要があります。後者なら、本人自身が持っているメンタルモデル(Mental Models)を問い直すことが、次の学習になるかもしれません。

したがって、Mindで不安が見えたからといって、その原因をすぐに決めてはいけません。**心理状態(Mind)は原因ではありません。原因を探すためのシグナルです。**

「自分の成果」とは何か

AIは、さらに佐藤さんへ聞きます。

「佐藤さんにとって、『自分が成果を出した』と感じるのは、どのようなときですか。」

佐藤さんは、「自分で顧客の問題を解決して、顧客から感謝されたときです」と答えました。

するとAIは、「では、佐藤さんが育てた現地社員が、自分で顧客の問題を解決し、顧客から評価された場合、それは佐藤さんの成果には含まれないのでしょうか」と問い返します。

佐藤さんは少し考え、「管理職になるなら、それも自分の成果と考える必要があるのかもしれませんが」と答えます。

会社から「管理職とはこういうものだ」と正解を教えられたのではありません。自分自身の経験と現在の Role を比べることで、「成果」という言葉の意味を本人が問い直しています。

第8章で設定した「**自分ができる**」から「**人を通じてできる**」へという成長テーマが、ここでは行動だけではなく、本人自身の「**成果の見方**」にまで入ってきたことになります。

Mind の対話から、メンタルモデル(Mental Models)が見えることがある

心理状態(Mind)の対話を掘り下げていくと、その背後に本人のメンタルモデル(Mental Models)が見えてくることがあります。

例えば、「自分でやった仕事でなければ、自分の成果ではない」「管理職なら、部下より多くの答えを持っていなければならない」「忙しいことは、会社から必要とされている証拠だ」といった前提です。

こうした考えは、本人にとってあまりにも自然なため、自分では「前提」と気づいていないことがあります。

しかし、「なぜそれが不安なのか」「なぜそれを達成感と感じるのか」と対話を続けると、自分がどのようなレンズを通して仕事を見ているのかが、少しずつ言葉になることがあります。

Mind を見る目的は、その考えを AI や会社が矯正することではありません。

本人が自分の見方に気づき、それが現在の Role や目指す未来に合っているのかを、自分で問い直せるようにすることです。

AI は、一度整理して本人へ返す

節目の振り返りでは、本人は多くのことを話します。達成感、不満、不安、負担、上司への考え、会社への期待、将来やりたいこと。対話が続けば、本人自身も、自分が何を話したのか整理しにくくなります。

そこで AI は、本人が話した内容を一度構造化し、鏡として本人へ返します。

例えば、「今回の振り返りでは、現地社員が自分で判断する場面が増えたことに強い達成感がある。一方、自分自身の直接案件が減ったことで、会社からどう評価されるのか不安がある。そして、海外拠点責任者になるためには、自分一人の成果だけではなく、チーム全体の成果を見る必要があると考え始めているように見えます。この整理は、佐藤さんの認識と合っていますか」と返します。

それに対して佐藤さんは、「二つ目は少し違います。案件が減ったことが嫌なのではなく、自分が何を基準に評価されるのか分からないことが不安なのだと思います」と訂正するかもしれません。

この訂正が重要です。AI の要約を正解にするものではありません。**AI が仮に構造化し、本人が確認し、必要なら修正する。**

心理状態(Mind)では、本人自身による意味づけを優先します。

過去の自分と比べると、変化が見える

Mind でも、AI は時間軸を持った鏡として使うことができます。

EvaliA を始めたころ、佐藤さんは、「自分で難しい案件を解決したときに、一番やりがいを感じます」と話していたとします。

半年後には、「現地社員が自分で解決できるようになったことが、一番うれしかった」と話しています。AI は、その違いを本人へ返します。

「半年前には、自分自身が難しい案件を解決したときに強いやりがいを感じると話していました。今回は、現地社員が自分で判断できるようになったことに最も達成感を感じています。この変化をどう考えますか。」

佐藤さんは、「自分が何を成果だと思うかが変わってきたのかもしれない」と気づくかもしれません。

日々の仕事の中では、自分自身の変化は見えにくいものです。しかし、過去の自分と現在の自分を比較することで、本人自身が自分の変化を認識できます。

エンゲージメント(Engagement)とは役割が違う

一般的なエンゲージメント・サーベイ(Engagement Survey)には、組織全体の状態を広く把握するという重要な役割があります。

一方、EvaliA で心理状態(Mind)を見る目的は、組織の平均点をつくることではありません。

本人が、「自分は今、この仕事をどう受け止めているのか」「なぜそう感じているのか」「何を成長と感じているのか」「何を变えたいのか」を、自分自身の言葉で整理することです。

したがって、本書では、エンゲージメント(Engagement)が組織や人の状態を見る一つの方法であるのに対して、EvaliA の心理状態(Mind)は、**本人がその状態の意味を考え、次の学習へつなげるための観測経路**として位置づけます。

本人で変えられることと、会社が変わることを分ける

佐藤さんの「自分がどう評価されるのか分からない」という不安について、AIはさらに、「その不安について、佐藤さん自身で確認できること、変えられることはありますか」と聞きます。

佐藤さんは、「上司と、自分の Role に対する評価基準を確認したいです」と答えました。

これは本人自身ができる行動です。例えば次の週次行動(Weekly Action)として、「上司との 1on1 で、自分に期待されている成果が個人案件だけなのか、チーム育成まで含むのかを確認する」と設定できます。

一方、その結果、本当に「管理職には部下育成を求めているのに、評価制度では個人売上しか見ていない」と分かったなら、本人の考え方を変えても解決しません。会社側の制度を見直す必要があります。

Mind から始まった一人の不安が、組織側の矛盾を発見する入口になることもあるのです。

同じシグナルが複数人に出たら、組織の問いになる

佐藤さん一人だけが、「部下育成を求められているのに、自分の売上だけで評価されているように感じる」と話しているなら、本人と上司の認識のずれかもしれません。

しかし、複数の管理職から同じシグナルが出ているのであれば、意味が変わります。

会社が期待している Role と、実際の評価基準は整合しているのか。部下育成を求めながら、そのために必要な時間を確保できる制度になっているのか。こうした構造を確認する必要があります。

必要であれば、SCAD へつなぎます。

不満や不安は、会社にとって単に消すべきノイズとは限りません。**組織の仕組みを見直すための学習情報になる場合があります。**

生の対話と、組織へ返す情報は分ける

心理状態(Mind)を扱うときには、情報の境界が特に重要です。

例えば佐藤さんが AI との対話の中で、「上司は自分の仕事を十分理解していないと思います」と話したとします。その文章がそのまま上司へ渡るのであれば、次回から佐藤さんは本当に考えていることを話しにくくなるでしょう。

本人の内省のための対話と、会社が学習へ使う情報は分ける必要があります。

本人自身が上司へ伝えたいのであれば、「自分の Role と評価基準を確認したい」という形で本人から伝えることができます。一方、組織レベルで扱うのであれば、「Role 期待と評価基準の認識にずれを感じているケースがある」と構造化できます。

さらに複数人に共通しているなら、「管理職の Role 期待と評価制度の整合性を確認する必要がある」という組織課題へ変えることもできます。

本人の生の言葉を公開するのではなく、そこから組織が学ぶ必要のある構造を取り出す。この境界がなければ、内省支援は監視へ変わってしまいます。

Mind を、人事評価の点数にしない

心理状態(Mind)から多くの情報が得られるようになると、それを数値化し、人事評価へ使いたくなるかもしれません。

しかし、「会社に満足している社員ほど高く評価する」「不満を持つ社員は低く評価する」という使い方をすれば、社員はすぐに本音を言わなくなります。

また、不満を持っている人が、会社への関心を失っているとは限りません。「この制度はおかしいのではないか」と考えるのは、むしろ会社を良くしたいからかもしれません。

心理状態(Mind)は、本人をランクづけするための情報ではありません。

本人と組織が、次に何を学ぶ必要があるのかを考えるための情報です。

半年前に描いた未来も、見直してよい

節目の振り返りでは、第8章で設定した本人の未来そのものも見直します。

佐藤さんは最初、「海外拠点責任者になりたい」と考えていました。半年間、実際に人へ仕事を任せ、チームを動かす経験をした結果、「やはり人を育てながら事業をつくる仕事をしたい」と、以前より強く思うかもしれません。

反対に、「実際にやってみると、自分は組織マネジメントより、専門家として難しい問題を解く方がやりたい」と気づくこともあるでしょう。それも失敗ではありません。

最初の Plan は、本人を縛る契約ではなく、学習を始めるための仮説でした。実際に行動し、経験したからこそ、自分の未来像そのものを選び直せます。

自己マスタリー(Personal Mastery)とは、一度決めた未来へ自分を固定することではありません。**現実から学びながら、ありたい姿そのものも問い直し続けることです。**

Mind は、本人にしか見えない情報を言葉にする

行動過程(Process)には、外からも比較的に見えやすい情報があります。何をしたのか。どこで止まったのか。どの行動が増えたのか。

一方、心理状態(Mind)には、本人が言葉にしなければ見えない情報があります。なぜうれしかったのか。なぜ不安なのか。何を成長と感じているのか。何に違和感があるのか。自分の未来をどう考えているのか。

AI が本人の心を読むものではありません。本人が、自分でも十分に言葉にできていなかったものを、対話によって言語化する。

AI は、そのための問いを返し、内容を整理し、過去との違いを鏡として返します。そして、**最終的にその意味を決めるのは本人です。**

二つの窓がそろった。しかし、まだ一つ足りない

ここまでで、佐藤さんについて二つの異なる情報が見えてきました。

行動過程(Process)では、佐藤さんは前進しています。現地社員へ意見を聞き、判断を任せ、自分がすべて答えを出さなくても仕事が進む場面が増えています。

心理状態(Mind)では、その変化に達成感を感じる一方で、「自分の直接案件が減り、自分は会社からどう評価されるのか」という不安も持っています。

行動だけを見れば、成長している。本人の受け止め方を見ると、成長に伴って新しい不安も生まれている。どちらも、佐藤さんの現在地です。

しかし、まだもう一つ重要な視点が残っています。

会社が期待していた成果は、本当に出ているのでしょうか。

本人は成長を感じている。行動も変わっている。それでも、会社として必要な成果につながっていない可能性があります。反対に、本人は自分に自信を持てていなくても、上司から見ると大きな成果を生み出しているかもしれません。

次章では、第三の観測経路である**成果(Result)**へ進みます。

重要業績評価指標(KPI)と Role に照らして、上司は佐藤さんの成果をどう見ているのか。ただし、人工知能(AI)が佐藤さんを採点するものではありません。

AIが上司へ問いを返し、評価する側自身も学習する。そこから、三つ目の窓が開きます。

第12章 上司は、何を見て成果を確認するのか

——重要業績評価指標(KPI)と成果(Result)を、AIとの対話で見る

第9章では、佐藤健太さんが毎週何を試し、どこで止まり、何を学びながら行動を変えてきたのかを、行動過程(Process)として見てきました。第11章では、佐藤さん自身がその経験をどのように受け止め、何に達成感を感じ、何に不安を持っているのかを、心理状態(Mind)という別の経路から振り返りました。

ここまでで、本人側から見た二つの情報があります。しかし、会社で仕事をしている以上、もう一つ確認しなければならないものがあります。

実際に、会社が期待した成果は出ているのか。

それが、三つ目の観測経路である成果(Result)です。

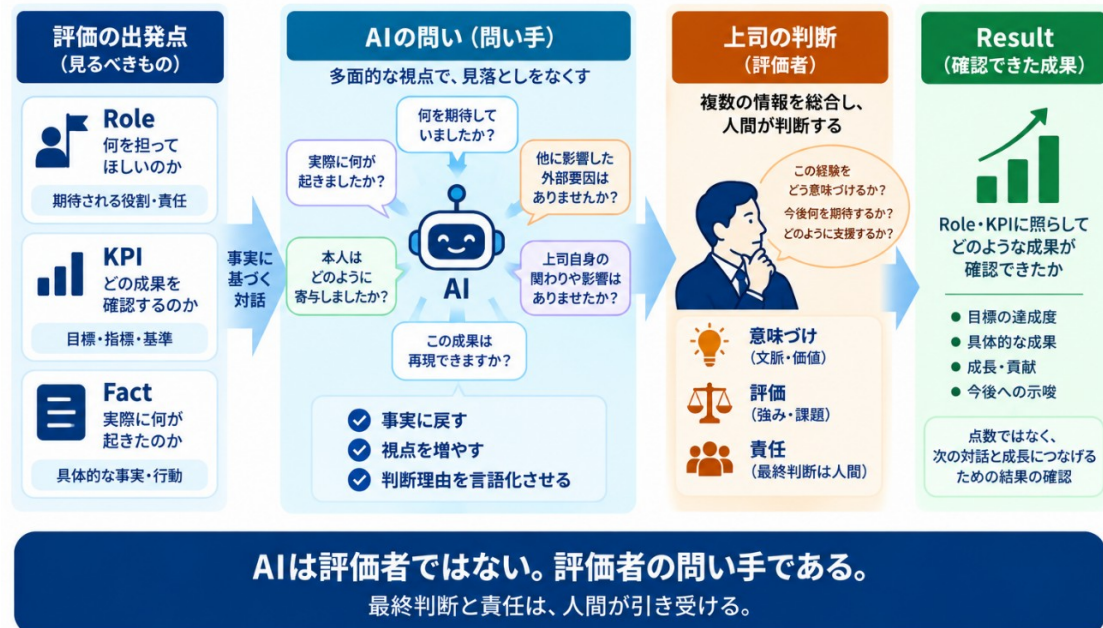
ここで大切なのは、人工知能(AI)が佐藤さんを評価するわけではないということです。AIが行うのは、上司へ問いを返し、上司自身が何を根拠に成果を見ているのかを具体化することです。

言い換えれば、**AIが社員を評価するものではありません。AIが、評価する上司へ問いを返すのです。**

図 12-1

AIは評価者ではない。評価者の問い手である

— Role・KPI・事実を、上司の判断へつなぐ —



最初から三つの情報を混ぜない

成果(Result)を確認するときには、一つ重要な原則があります。行動過程(Process)、心理状態(Mind)、成果(Result)を、最初から一つに混ぜないことです。

例えば上司が佐藤さんの成果を確認する前に、「佐藤さんは最近かなり努力しています」「本人は評価されるか不安を感じています」といった Process や Mind の詳細をすべて見ていたらどうなるでしょうか。

上司は無意識にその情報へ引っ張られる可能性があります。「これだけ努力しているのだから、高く評価してあげよう」と考えることもあれば、「本人が不安を感じているなら、まだ自信が足りないのだろう」と見ることもあるでしょう。

しかし、それでは成果(Result)という第三の経路を別に設ける意味が薄れます。

だから成果(Result)を確認するときには、まず、本人への Role を期待していたのか、どの重要業績評価指標(Key Performance Indicator)を設定していたのか、実際に何が起きたのか、本人はその成果への程度寄与したのかを、できるだけ成果側の事実から見ます。

三つを比較するのは、その後です。別々に観測するからこそ、後から見つかる Gap に意味が生まれます。

成果(Result)は、人事評価そのものではない

もう一つ、区別しておきたいことがあります。本書でいう成果(Result)は、昇給、賞与、昇格などを最終的に決める人事評価そのものではありません。

成果(Result)とは、会社が本人へ期待した Role と KPI に対して、**上司という別の視点から成果を確認する観測経路**です。

もちろん、その情報が最終的な人事評価の材料になることはあります。しかし、Process、Mind、Result へそれぞれ点数をつけ、それを合計して「総合人材スコア」をつくることが目的ではありません。

それをしてしまえば、三つを別々に見る意味がなくなります。

EvaliA で重要なのは、「三つの合計はいくつか」ではありません。**三つが同じ方向を示しているのか。それとも、違うものを示しているのか**を見ることです。

半年前、佐藤さんへ何を期待していたのか

佐藤さん自身が目指していたのは、「現地社員を育て、自分がいなくてもチームとして成果を出せる経営人材になる」という未来でした。

会社としても、海外事業を拡大し、現地社員が自律して顧客対応を行える状態をつくりたい。その接点から、上司は佐藤さんへ、「自分で顧客対応を完結するだけでなく、現地社員へ判断を移し、本人たちだけでも一定範囲の仕事を完結できる状態をつくる」という Role を期待しました。

そして、その Role が実際の成果へつながっているかを見るために、現地社員だけで完結できた案件の割合、佐藤さんが直接介入せずに処理できた業務範囲、顧客対応品質などを KPI として設定したとします。

半年（または四半期）が経過しました。上司は、この Role と KPI に照らし、佐藤さんの成果(Result)を確認します。

いきなり「何点ですか」と聞かない

従来の評価制度であれば、「佐藤さんを 5 段階で評価してください」と聞くかもしれません。

しかし、いきなり点数を求めると、上司の第一印象がそのまま評価になることがあります。「佐藤さんは優秀だ」「最近は少し頼りない」「以前から何でも自分でやってしまう」。そうした印象を先に持ち、その後で、その印象に合う事実だけを集めてしまうこともあります。

そこで AI は、先に事実へ戻すための問いを返します。

「この半年、佐藤さんへ期待していた主な Role は何でしたか。」「その Role について、実際にどのような変化が確認できましたか。」「設定していた KPI は、どこまで達成されましたか。」「その成果の中で、佐藤さん本人が寄与したと考えられる部分はどこですか。」

上司が、「現地社員だけで顧客対応を完結できる案件が増えました」と答えたとします。AI はさらに、「その変化に、佐藤さんはどのように関わったのでしょうか」と問い返します。

上司は、「以前は佐藤さん自身が最終回答までつくっていました。最近は、現地社員が先に案を出し、佐藤さんは必要な部分だけを支援する形に変わっています」と答えるかもしれません。

こうして、評価を先に置くのではなく、事実を先に言葉にするのです。

Role と KPI は、両方を見る

成果(Result)では、KPI だけを見ればよいわけではありません。

例えば、現地社員だけで完結した案件数が増えたとします。数字だけを見れば成果です。しかし実際には、佐藤さんが裏でほとんど答えを教え、最後だけ現地社員に処理させているのであれば、「現地社員の自律をつくる」という Role は十分に果たされていません。

反対に、KPI の数字はまだ大きく伸びていなくても、以前はすべて佐藤さんへ確認していた現地社員が、自分で案を考えたうえで相談するようになってい
るなら、Role として重要な変化が起きている可能性があります。

KPI は、何が起きたのかを見るために必要です。一方、Role は、**本人が何を担う人として、その成果をつくったのか**を見るために必要です。

KPI だけなら、数字を出した人だけが評価されやすくなります。Role だけなら、「よく頑張りました」という主観へ流れやすくなります。これら両方を見ることが重要です。

佐藤さん個人の数字だけを見ると、悪化している

佐藤さんには、興味深い変化が起きています。

現地社員へ仕事を任せるようになったため、佐藤さん自身が直接処理する案件数は減りました。本人の案件数だけを見れば、「以前より生産性が落ちた」と評価されるかもしれません。

しかし同じ期間に、現地社員三人が自分たちで処理できる案件は増えました。佐藤さん一人の処理量は減った一方で、チーム全体の処理能力は上がっています。

ここで AI は、上司へ問いを返します。

「佐藤さん本人の直接処理件数は減っています。一方、チーム全体で処理できる案件数は増えています。今回設定した Role に照らすと、この変化をどのように見ますか。」

この問いによって、上司自身も「成果とは何か」を考えることになります。プレーヤーとして本人が一番多く仕事をすることを成果とするのか。それとも管理職候補として、**他者を通じて、より大きな成果を出せる状態をつくること**を成果とするのか。

Role が変わったのに、評価する数字だけが以前のままであれば、会社は本人へ矛盾したメッセージを送ることになります。

成果が出たときにも、「なぜ」を聞く

成果(Result)の対話は、未達の原因を探すためだけにあるわけではありません。

良い成果が出たときにも、AI は問いを返します。

なぜ今回の成果が出たのか。佐藤さん本人が寄与した部分はどこか。現地社員自身の成長による部分はどこか。仕組みの変更や市場環境など、本人以外の要因はなかったか。そして、その成果は次の半年にも再現できるのか。

例えば売上が大きく増えていたとしても、大口顧客から偶然大型案件が入っただけなら、その結果をそのまま本人の能力向上とは言えません。

一方、現地社員へ判断基準を明確にし、顧客対応方法を標準化した結果として、チーム全体の処理能力が継続的に上がっているのであれば、再現性のある成果である可能性があります。

良い結果が出たときほど、「なぜ出たのか」を言葉にすることが重要です。その理由を説明できれば、後に他の人やチームへ共有できる知識になるからです。

未達でも、すぐ本人の能力不足とは考えない

反対に、KPI が未達だった場合はどうでしょうか。

例えば佐藤さんのチームが、新規顧客獲得目標を達成できなかったとします。上司がすぐに、「佐藤さんにはまだ営業力が足りない」と考えたとします。

AIはそこで、「未達の理由をどのように考えていますか」と問い、「佐藤さん本人がコントロールできた要因は何でしたか」「本人では変えられなかった要因は何ですか」「商品、価格、市場、他部署、上司の支援などの影響はありませんでしたか」と視点を広げます。

これは、本人の責任をなくすためではありません。結果には、多くの要因が混ざっているからです。

市場が悪かったとしても、本人が変化を早く察知し、営業方法を変えられた可能性はあります。商品の問題に気づいたなら、上司やSCADへ共有できたかもしれません。現地社員へ任せられなかったのであれば、Roleや育成方法を見直すこともできたかもしれません。

外部要因を考慮しながらも、その環境の中で本人が何を変えられたのかを見る。結果だけでも、言い訳だけでもない。その間を見る必要があります。

本人の寄与と、環境の影響を分けて考える

成果(Result)の難しさは、結果を完全に一人の人間へ帰属できないことです。

売上が大きく伸びたとしても、市場そのものが成長していたのかもしれませんが。競合が撤退した可能性もあります。会社が新しい商品を投入したことが大きかったかもしれません。

反対に、売上が未達でも、本人の行動そのものは適切だった場合があります。

そこでAIは、「今回の成果のうち、本人の寄与と考えられるものは何ですか」「本人以外の要因には何がありますか」と上司へ問い返します。

完全に切り分けることはできません。それでも、一度分けて考えるだけで、「結果が良いから本人も優秀」「結果が悪いから本人の能力不足」という短絡を減らすことができます。

上司自身も、成果をつくる環境の一部である

部下を評価するときに忘れられやすいのが、上司自身の行動です。

例えば上司は佐藤さんへ、「現地社員へもっと判断を任せてください」という Role を与えていたとします。ところが、重要な顧客案件になるたびに、上司が佐藤さんを飛び越えて現地社員へ直接指示していた。

これでは、佐藤さんが現地社員へ権限を渡すことは難しくなります。

そこで AI は、上司へ「佐藤さんがこの Role を果たすために、上司としてどのような支援をしましたか」「反対に、佐藤さんが Role を果たしにくくなるような関わりはありませんでしたか」と問い返すことができます。

これは上司を責めるためではありません。**上司自身も、成果が生まれるシステムの一部だからです。**

部下の Result を見るときに、本人だけではなく周囲の構造も見る。ここでも、システム思考(Systems Thinking)が実務として働きます。

上司にも、メンタルモデル(Mental Models)がある

評価する上司は、中立で完全に客観的な存在ではありません。

上司にも過去の経験があり、「管理職とはこういうものだ」「この社員はこういう人だ」という見方があります。つまり、上司にもメンタルモデル(Mental Models)があります。

例えば上司が以前から、「佐藤さんは何でも自分でやってしまう」という印象を持っていたとします。その印象が強ければ、佐藤さんが十回部下へ任せられるようになっていても、一度だけ自分で全部処理した場面を見て、「やはり佐藤さんは変わっていない」と判断する可能性があります。

反対に、「佐藤さんは優秀だから」という印象が強ければ、同じ失敗を好意的に解釈することもあるでしょう。

そこで AI は、「その評価をした具体的な行動は何ですか」「反対の事実はありませんでしたか」「半年前と現在で変わったものは何ですか」と問い返します。

AIが上司の評価を否定するのではありません。印象を、具体的な事実へ戻すのです。

EvaliAが育てるのは、被評価者だけではない

ここから、EvaliAのもう一つの役割が見えてきます。評価される社員だけが学ぶのではありません。評価する「上司」もまた、学習主体です。

人を評価することは難しい仕事です。数字だけを見れば短期成果を出す人が有利になります。努力だけを見れば成果責任が曖昧になります。印象へ頼れば評価が属人的になります。

そのたびにAIから、「どの事実からそう判断しましたか」「本人の寄与はどこですか」「外部要因はありませんでしたか」「上司としての支援は十分でしたか」「その成果は再現できますか」と問われる。

こうした問いを繰り返すことで、上司自身の評価の仕方にも学習が起こります。EvaliAが育成するのは、被評価者だけではありません。評価する上司もまた、学習するのです。

AIへ評価を任せれば、公平になるのか

ここまでAIが上司へ多くの問いを返せるのであれば、「最終的な評価もAIに任せればよいのではないか」と考えるかもしれません。

しかし、本書ではそう考えません。評価には価値判断が含まれるからです。

今期の成果をどの程度重視するのか。将来の成長可能性をどう見るのか。新しい仕事へ挑戦して失敗した人をどう見るのか。本人の直接成果と、チームを通じた成果をどう考えるのか。

これらは、会社の戦略や価値観によって変わります。AIにできるのは、評価材料を整理し、見落としている可能性のある視点を示し、上司が自分の判断理由を具体化するための問いを返すことです。

AIは評価者ではありません。評価者の問い手です。最終的な判断と責任は、人間が持ちます。

点数を入力する前に、問いへ答える

成果(Result)を確認するとき、固定された評価項目を大量に並べ、上司が順番に点数を入力するだけでは十分ではありません。

佐藤さんがKPIを大きく達成しているなら、「なぜ達成できたのか」「来期も再現できるのか」「本人だけに依存した成果ではないか」という問いが重要になります。

一方、未達であれば、「どこで計画との差が生まれたのか」「本人が変えられたものは何か」「外部要因は何だったのか」「同じ条件でも次回変えられることは何か」という問いが必要になります。

また、佐藤さんのような管理職候補であれば、「本人自身の成果だけではなく、チーム全体の能力は上がったのか」という問いも重要です。

人によって、RoleもContextも違います。

だからAIが状況に応じて問いを変えることで、成果確認を単なる点数入力から、**成果の意味を考える対話**へ変えることができます。

良い成果を、組織学習の入口にする

佐藤さんのチーム全体の能力が上がったとします。

ここで「良い評価だった」で終われば、それは佐藤さん個人の成果です。しかし、「なぜ佐藤さんは現地社員へ任せられるようになったのか」「その方法の中で、他の管理職にも使えるものは何か」と考えれば、話は組織学習へ進みます。

AIは、「今回の成果を次の管理職でも再現するために、残せる知識はありますか」と上司へ問い返すことができます。

もし答えが、「佐藤さんだからできた」で終われば、組織としてはまだ十分に学べていません。

何を見ていたのか。どんな問いを部下へ返していたのか。どの判断を本人へ任せ、どこだけ自分が持っていたのか。

そこまで言葉にできれば、一人の成功から次の人が学ぶことができます。

この流れは、第 14 章で扱う名人循環(Mastery Loop)へつながっていきます。

未達から、会社の設計問題が見えることもある

反対に、一人だけではなく、複数人が同じ KPI を未達にしているとします。

上司への成果(Result)対話を重ねる中で、「価格が市場に合っていない」「承認が遅く、顧客を逃している」「新規開拓を求められているのに、既存顧客対応だけで時間を使い切っている」といった共通する理由が見えてきた。

この場合、社員一人ひとりへ研修を追加しても、本質的な問題は解決しないかもしれません。

商品、価格、Role、KPI、業務プロセス、会社の仕組みそのものを見る必要があります。必要であれば、SCAD へ上げます。

成果(Result)の確認は、人をランクづけするためだけの作業ではありません。**会社の設計そのものに矛盾がないかを発見する入口にもなるのです。**

佐藤さんを、上司はどう見たのか

では、半年後の佐藤さんを、上司はどう見たのでしょうか。

佐藤さん本人が直接処理する案件は減っています。一方で、現地社員だけで完結できる案件は増えました。顧客対応品質も大きく落ちていません。佐藤さんが不在でも進められる仕事が増えています。

上司は、例えばこう考えていたとします。「以前は、佐藤さん自身が一番仕事のできる人でした。今は、周囲の人を通じてチームとして成果を出す方向へ変わってきています。重要案件ではまだ自分で抱える傾向がありますが、半年前より明らかに変化しています。」

これは、第 11 章で佐藤さん自身が感じていた不安とは少し違います。

佐藤さんは、「自分の案件数が減ったので、会社からの評価も下がるのではないか」と考えていました。しかし上司は、「本人の案件数が減ったことよりも、チーム全体の能力が上がったことを評価している」と考えていた。

ここに、初めて明確な Gap が現れます。

本人と上司の見方が違ってもよい

本人と上司の見方が違うと、「どちらが正しいのか」と考えたくなります。

しかし、必ずしもどちらかが間違っているわけではありません。

佐藤さんは、自分自身の経験、達成感、不安を見ています。一方、上司は、会社が期待した Role と成果を見えています。見ている窓が違うのです。

本人は成長していると感じているのに Result が出ていないこともあります。Result は高いのに、本人は仕事に意味を感じていないこともあります。本人の行動は変わっているのに、上司がその変化にまだ気づいていないこともあります。

重要なのは、最初から一致させることではありません。**違っているなら、なぜ違うのかを考えることです。**

三つの観測経路が、ここでそろった

ここまでで、佐藤さんについて三つの観測経路がそろいました。

行動過程(Process)では、佐藤さんが何を試し、どこで止まり、どのように行動を変えてきたのかを見えています。心理状態(Mind)では、佐藤さん自身がその

経験をどのように受け止め、何に達成感を感じ、何に不安を持っているのかを見ています。そして成果(Result)では、会社が期待した Role と KPI に対して、実際にどのような成果が生まれたのかを、上司という別の視点から確認しました。

佐藤さんの場合、Process は前進しています。Mind には、現地社員の成長への達成感と同時に、評価への不安があります。そして Result も前進しています。

ここで三つを一つの総合点へまとめる必要はありません。

むしろ価値があるのは、「**なぜ Process と Result は前進しているのに、Mind には評価への不安が残っているのか。**」という違いです。

本人のメンタルモデル(Mental Models)を問い直すべきなのか。上司と評価基準を確認すべきなのか。それとも、会社の Role 設計や評価制度そのものに矛盾があるのか。

次章では、ここまで意図的に別々に取ってきた**行動過程(Process)**、**心理状態(Mind)**、**成果(Result)**を初めて比較します。

目的は、人を一つの数字で評価することではありません。**三つの経路の間にある Gap から、「次にどこを変えるべきか」を見つけることです。**

第13章 三つの経路を比較すると、何が見えるのか

——行動過程(Process)・心理状態(Mind)・成果(Result)の Gap から、次に変える場所を見つける

第9章では、佐藤健太さんが毎週どのような行動を試し、どこで止まり、何を学んだのかを、行動過程(Process)として見てきました。第11章では、佐藤さん自身がその経験をどのように受け止めているのかを、心理状態(Mind)という本人側の経路から振り返りました。そして第12章では、本人の詳細な Process や Mind にできるだけ引っ張られない形で、上司が役割(Role)と重要業績評価指標(Key Performance Indicator)に照らして成果(Result)を確認しました。

三つは、最初から混ぜませんでした。なぜなら、EvaliA で本当に見たいのは、Process、Mind、Result の合計点ではないからです。

三つの経路が、どこで一致し、どこで違っているのか。その違いの中に、次の学習への手がかりがあります。

佐藤さんには、三つの異なる姿が見えていた

佐藤さんの半年間を、もう一度三つの窓から見てみましょう。

行動過程(Process)では、明らかな前進がありました。以前は難しい案件になるほど自分で答えを出していましたが、最近は現地社員へ先に案を出してもらい、その理由を聞き、必要なところだけを支援する場面が増えています。

成果(Result)にも変化がありました。佐藤さん自身が直接処理する案件数は減りましたが、現地社員だけで完結できる案件は増え、チーム全体の処理能力も上がっています。上司も、「本人が一番多く仕事をする状態から、人を通じて成果を出す状態へ変わり始めている」と見ています。

一方、心理状態(Mind)では、別の情報が見えていました。佐藤さんは、現地社員が成長していることには達成感を感じています。しかし、自分自身が直接処理する案件が減ったことで、「自分は会社からどう評価されるのだろう」という不安も持っています。

つまり、Process は前進している。Result も前進している。しかし Mind には、達成感と同時に評価への不安が残っています。

どれか一つだけを見れば、違う判断になります。Process だけなら「順調に成長している」と見えるでしょう。Result だけなら「管理職候補として良い成果が出ている」と見えます。Mind だけなら、「本人に不安があるので、何か問題が起きている」と考えるかもしれません。

しかし三つを並べると、別の問いが生まれます。「**行動も成果も前進しているのに、なぜ本人には評価への不安が残っているのか。**」

この違いが、Gap です。

Gap は、問題判定ではなく「探索地点」である

Gap という言葉を使うと、「ずれているから悪い」「何か異常がある」と感じるかもしれません。しかし、本書でいう Gap は、そのような意味ではありません。

人が変わろうとしている途中では、むしろ Gap が生まれることは自然です。

佐藤さんは長い間、プレーヤーとして成功してきました。難しい仕事を自分で解決し、顧客から直接感謝される。それが本人にとっての「成果」でした。

ところが今は、管理職候補として、人へ判断を渡す方向へ行動が変わっています。実際の成果も、本人の直接成果から、チームを通じた成果へ変わり始めています。

行動は先に変わった。成果も変わった。しかし、「自分の成果とは何か」という本人の認識は、まだ完全には追いついていない。そう考えることもできます。

したがって、Gap は失敗を意味するものではありません。

「ここには、もう少し詳しく見る価値がある」という探索地点です。

Gap があるから問題なのではなく、Gap から何を問い、どこを確かめるかが重要なのです。

佐藤さんの場合、どこを変えるべきなのか

佐藤さんに Gap が見つかったからといって、すぐに「本人をもっと教育しよう」と考える必要はありません。

AI は、例えば佐藤さんへ、「以前は、自分で案件を解決することを成果だと考えていました。現在は、佐藤さん自身の案件数は減っていますが、現地社員だけで完結する案件は増えています。現在の Role から見ると、この二つをどのように考えますか」と問い返すことができます。

これは、本人のメンタルモデル(Mental Models)を問い直す方向の介入です。

しかし、佐藤さんがさらに、「上司は部下育成を評価すると言っていますが、人事評価表を見ると個人売上しかありません」と話したのであれば、話は変わります。

その場合、本人の考え方を変える前に、会社が期待する Role と評価制度が本当に整合しているのかを確認する必要があります。

同じ「評価への不安」という Mind のシグナルでも、本人のメンタルモデルに主な原因があるのか、上司との認識のずれなのか、それとも評価制度そのものの矛盾なのかによって、変える場所は異なります。

AI は Gap を見つけても原因を確定しません。**Gap は原因ではなく、原因を探す場所なのです。**

AI が出すのは、原因ではなく原因仮説である

三つの情報がそろえば、AI は多くの推測をできるようになります。

しかし、「佐藤さんの問題は自己評価の低さです」「会社の評価制度に問題があります」と断定してはいけません。

AIができるのは、「本人の成果認識と現在の Role との間にずれがある可能性があります」「本人と上司で評価基準の認識を確認する価値があります」「期待 Role と評価制度が整合しているか確認する必要があるかもしれません」といった、複数の原因仮説を返すことです。

そこから人間が、実際の事実を確認します。AI が原因を決めるのではありません。人間が原因を探するための視野を広げることが、AI の役割です。

同じ「社員を見る」でも、Gap の組み合わせで意味は変わる

佐藤さんの場合は、Process と Result が前進し、Mind に不安が残っていました。

しかし、実際の組織ではさまざまな組み合わせが起こります。ここからは、佐藤さんとは異なる三人の説明用ペルソナを見てみましょう。いずれも特定の実在社員ではなく、三つの経路を比較したときに何が見えるのかを説明するためのケースです。

ケース 1 行動も意欲も前進しているのに、成果が出ない

中村美咲さん、32 歳。営業担当です。

中村さんはこの半年、新規顧客開拓に積極的に取り組んできました。週次 AI コーチングを見ると、新しい見込み客へ連絡し、提案内容を変え、断られた理由を振り返りながら改善しています。Process では、明らかな前進があります。

本人も、「以前より顧客へ質問できるようになりました。自分でも成長していると思います」と話しており、Mind も前向きです。

しかし Result を見ると、新規受注はほとんど増えていません。

この状態で、「もっと頑張りましょう」と本人へ返すだけなら、介入先を間違えている可能性があります。中村さんは、すでにかなり動いているからです。

ここでは、本人の営業方法だけではなく、対象としている顧客層は正しいのか、価格は市場に合っているのか、商品に競争力があるのか、提案後の承認が遅すぎないか、設定された KPI が本当に受注につながるものなのか、と視野を広げる必要があります。

特に、同じ営業チームの複数人に同じ Gap が出ているなら、本人をさらに教育する前に、商品、価格、戦略、業務プロセスといった構造側を見る必要があります。

結果が悪いから人を直す。これは最も簡単で、しかし最も間違えやすい結論です。

ケース 2 成果は出ている。しかし学習が止まり始めている

高橋誠さん、44 歳。ベテラン営業社員です。

高橋さんは今期、大きな売上を上げています。成果(Result)だけを見れば、申し分ありません。本人も「今期は順調です」と感じており、心理状態(Mind)にも大きな問題は見えません。

ところが行動過程(Process)を見ると、別のものが見えてきます。

売上の多くが長年取引している大口顧客から生まれており、新しい顧客への活動はほとんどありません。新しい提案を試す行動も以前より減っています。

今期だけを見れば、大きな問題はないかもしれませんが。しかし、その大口顧客との取引がなくなったらどうなるでしょうか。

現在の成果が良いことと、将来も同じ成果を再現できることは同じではありません。

そこで AI は上司へ、「現在の成果のうち、来期にも再現できる部分はどこですか」「今回の成果は本人の新しい行動によって生まれたものですか。それとも、過去から続く取引関係による部分が大きいですか」と問い返すことができます。

高橋さんを低く評価するためではありません。**現在の成功の中に、将来の停滞が隠れていないかを見るためです。**

ケース 3 成果は出ているが、本人は「このままでは続かない」と感じている

山田彩さん、38 歳。プロジェクトマネージャーです。

山田さんのプロジェクトは、納期も品質も目標を達成しています。Process を見ると、毎週多くの問題へ対応し、関係部署を動かし、プロジェクトを前へ進めています。Result も良好です。

しかし、節目の Mind では、「今の働き方を、あと半年続けるのは難しいと思います」と話しています。

ここで「成果が出ているから問題ない」と考えれば、Mind のシグナルを無視することになります。

反対に、「負担を感じているなら、すぐ仕事量を減らそう」と即断するのも早いでしょう。新しい Role へ挑戦している途中で一時的に負荷が高いだけかもしれません。本人自身が一定期間の負荷を受け入れている可能性もあります。

一方で、山田さんだけに判断や調整が集中し、本人がいなければプロジェクトが回らない構造になっている可能性もあります。

したがって、「何が負担なのか」「本人だけに集中している仕事は何か」「Role や権限を分散できないか」「チーム構成に問題はないか」を確認します。

Mind の悪化を診断結果として扱うものではありません。**現在の成果を、このまま持続できる構造になっているのかを問い直すシグナル**として扱います。

四つのケースを並べると、見えるものが変わる

ケース別の観測結果と主な探索方向

ケース	行動過程 (Process)	心理状態 (Mind)	成果 (Result)	主な探索方向
佐藤さん	前進	達成感＋評価への不安	前進	成果認識、上司との 評価基準、制度
中村さん	前進	前向き	未達	営業方法、商品、価格、 市場、KPI
高橋さん	停滞	良好	良好	再現性、将来リスク、 次の挑戦
山田さん	前進	強い負担感	良好	Role、業務集中、 チーム構造、持続性

※ここでのいう EvaliA の「前進」「良好」「未達」等は、各ケースの状態を説明するための要約表現です。Process、Mind、Result を共通尺度で点数化したものではありません。

四人を並べると、「Result が良いから問題ない」「Mind が悪いから問題だ」といった単純な判断では、人の状態を十分に捉えられないことが分かります。

ここで、本書で繰り返してきた原則が重要になります。一つの数字、一つの行動、一つの評価、一つの視点だけで、人を決めない。複数の経路を持ち、その違いから問いをつくる。

Gap を見つけたら、「誰を変えるか」ではなく「どこを変えるか」

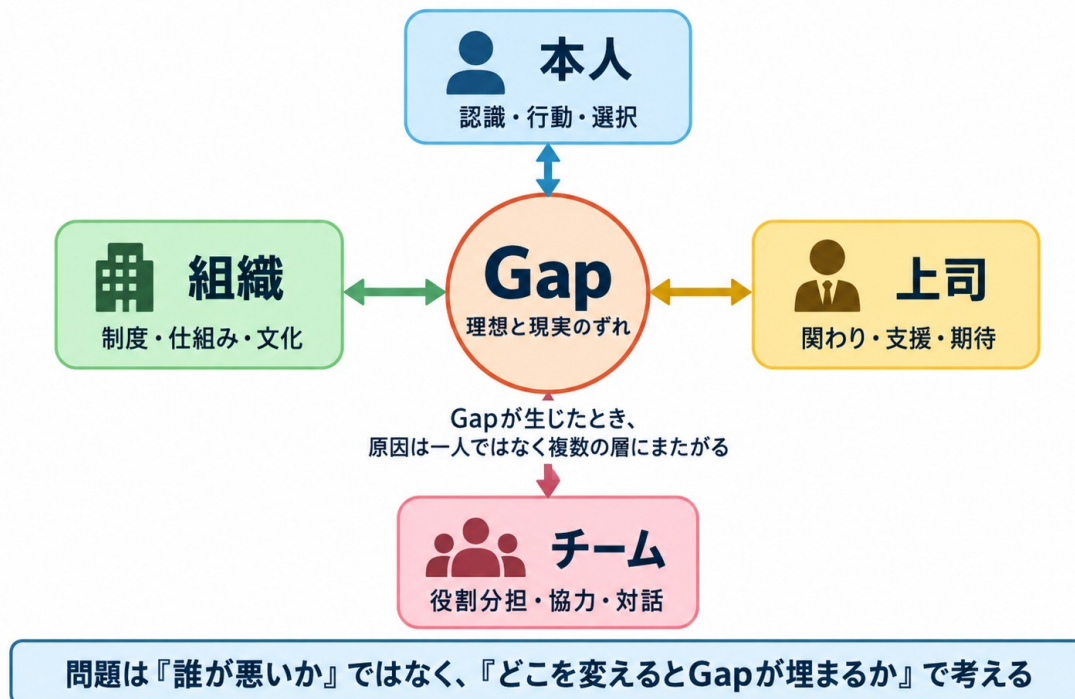
Gapが見つかり、会社はつい、「では、この社員をどう改善するか」と考えます。しかし、それではすべてが人材教育へ戻ってしまいます。

変えるべき場所は、大きく四つあります。本人、上司、チーム、組織です。

ただし、「まず本人を変え、それでだめなら上司へ」という順番ではありません。Gapの内容を見て、どこへ働きかけるのが適切なのかを考えます。

図13-1 変えるべきは『人』ではなく『場所』

13章のメッセージ：誰を変えるかではなく、どこを変えるか



本人へ返す——次の週次行動(Weekly Action)

本人自身で変えられることであれば、毎週の学習循環へ戻します。

例えば佐藤さんが、「自分が直接処理した仕事だけを成果として見ていた」と気づいたなら、次の Weekly Action として、部下が自分で完結した仕事を振り返り、自分の Role との関係を考えることができます。

中村さんの商談方法そのものに課題があると分かったなら、次の営業場面で一つ新しい方法を試すことができます。

本人で変えられるものは、本人の次の実験へ返します。

上司へ返す——Feedback、支援、権限、Role

本人では変えられなくても、上司との関係で変えられるものがあります。

佐藤さんへ「現地社員へ任せてほしい」と言いながら、重要案件では上司自身がすべて決めてしまうのであれば、本人へさらに「任せてください」と要求しても解決しません。

上司が、どこまで権限を渡すのか。どこから相談するのか。何を成果として見るのか。そこを明確にする必要があります。

この場合、介入先は本人教育ではなく、上司の Feedback、支援、権限、Role の調整です。

チームへ返す——Team CAGE と Team Action

複数人に同じ Gap がある場合には、チームへ戻します。

例えば営業チームの多くが、Process では積極的に新しい提案を出しているのに、Result へつながっていない。Team CAGE Map を見ると、探究 (Adventurer) の行動は多い一方で、成果 (Challenger) として担当者や期限を決め、最後までやり切る行動が不足していると分かったとします。

その場合、一人ひとりへ別々に「実行力を高めましょう」と教育するより、会議の進め方を変え、Role を明確にし、Team Action として「提案が出たら、必ず担当者と次の期限を決める」といった実験を行う方が適切かもしれません。

複数人の Gap は、チームの働き方そのものを学習対象にする入口になります。

組織へ返す——SCAD で仕組みを変える

さらに、チームだけでは変えられない問題もあります。

会社が管理職へ「部下を育ててほしい」と求めながら、人事評価では個人売上だけを評価しているのであれば、これは佐藤さん一人の問題ではありません。

営業社員が十分に行動しているにもかかわらず、会社の価格設定が市場に合わないため受注できないのであれば、それも本人の努力だけでは変えられません。

こうした問題は SCAD へ上げ、制度、KPI、標準、業務プロセス、システムなどを見直します。

人に求めている行動と、会社の仕組みが矛盾しているなら、人だけではなく仕組みを変える。ここに、個人学習と組織学習がつながります。

同じ Gap が複数人に現れると、共通 Gap になる

一人の Gap だけでは、本人固有の問題なのか、組織の問題なのか分からないことがあります。

しかし、例えば同じ部署の 10 人のうち 8 人について、Process では新しい行動が起き、Mind も前向きなのに、Result だけが未達という状態が続いていたらどうでしょうか。

全員の能力が不足している可能性もあります。しかし、商品、市場、KPI、承認プロセス、Role など、共通する構造を疑う必要もあります。

AI を組織学習へ使う意味はここにあります。

一人の中では個人課題に見えていた Gap が、複数人に共通すると、組織課題を探すシグナルへ変わる。

ただし、共通 Gap が見えたからといって、「これは制度の問題だ」と人工知能(AI)が原因を決めるわけではありません。

共通 Gap が示すのは、「個人だけではなく、構造側も調べる必要がある」ということです。

介入もまた、仮説である

原因を考え、どこを変えるか決めたとしても、その介入が正しいとは限りません。

例えば山田さんの負担を減らすために、仕事を他の社員へ分散した。その結果、Mind は改善したが、Result が大きく下がったとします。

あるいは、承認を速くするためにプロセスを簡略化した。Process は速くなったが、安全(Guardian)に関する確認が不足し、ミスが増えた。こうしたことは起こり得ます。

よって、何かを変えたら、もう一度観測します。Process はどう変わったのか。Mind はどう変わったのか。Result はどう変わったのか。新しい Gap は生まれていないか。

必要であれば、介入そのものを変えます。

介入(Intervention)もまた、正解ではなく仮説です。変えて、観測し、また変える。この繰り返しそのものが学習になります。

ここで、人材成長循環(Human Development Loop)が一周する

第8章からの流れを振り返ると、人材成長循環(Human Development Loop)がここで一周します。

計画(Plan)では、本人のありたい姿と会社の未来を接続し、Role と KPI を置きました。実行(Do)では、その Role を小さな週次行動(Weekly Action)へ落とし、実際の仕事の中で試しました。

確認(Check)では、行動過程(Process)、心理状態(Mind)、成果(Result)という三つの経路から状態を観測し、その間にある Gap を比較しました。そして改善(Act)では、本人、上司、チーム、組織のどこへ働きかけるのかを判断します。

その結果を、再び次の Plan へ戻します。

つまり、『**ありたい姿 → Role ・ KPI → Weekly Action → Process ・ Mind ・ Result → Gap → Intervention → 次の Plan**』という循環です。

これが、EvaliA によって人と組織の学習を回す、人材成長循環(Human Development Loop)です。

評価は、終点ではなく次の Plan への入力になる

従来の人事制度では、評価が一つの終点になりがちです。一年間働き、上司が評価し、昇給や賞与が決まり、そこで一度終了します。

しかし、学習する組織(Learning Organization)では、評価や振り返りを終点にはしません。

佐藤さんは半年間の経験から、「海外拠点責任者になるなら、自分が最も仕事のできる人になることより、自分がいなくても現地社員が成果を出せる状態をつくることの方が重要なのではないか」と考えるようになりました。

第 8 章で設定した「海外拠点責任者になりたい」という未来そのものが、少し深くなっています。

次の半年では、単に「部下へ任せる」だけではなく、「誰にどこまで判断を渡すのか」「どの判断を標準化するのか」「自分が不在でも仕事が進む Role 設計をどうつくるのか」という、より高い学習課題へ進むことができます。

人が学んだ分だけ、次の Plan で立てられる問いも高度になるのです。

三つの経路から、「この人はこういう人」と決めない

Process、Mind、Result を比較すると、多くの情報が得られます。

しかし、それを使って、「佐藤さんは管理職適性が高い」「中村さんは成果を出しにくいタイプだ」「山田さんは負荷に弱い」と人物像を固定してはいけません。

見えているのは、**現在の Role と Context における状態**です。

Role が変われば行動は変わります。上司が変われば本人の強みの出方も変わります。チームが変われば、それまで出ていなかった行動が現れることもあります。そして、本人自身も学習して変わります。

EvaliA が残すべきなのは、固定された人物ラベルではありません。**その人が、どのような状況で、何を試し、何に気づき、どう変わってきたのかという学習履歴**です。

学習履歴が増えるほど、次の問いは深くなる

佐藤さんが EvaliA を使い始めたころ、人工知能(AI)が返していた問いは、「部下へ先に考えてもらうことはできますか」というものでした。

半年後、佐藤さんはすでに一定程度、部下へ判断を渡せるようになっていきます。その人へ、いつまでも同じ質問をしても意味がありません。

次には、「どの仕事を誰へ任せるのか」「どこまで権限を渡すのか」「どの判断を標準化できるのか」「自分がいなくても判断できる仕組みをどうつくるのか」といった、より高い問いへ進みます。

過去の学習履歴があるからこそ、人工知能(AI)は、本人の以前と現在を比較し、次に必要な問いを変えることができます。

しかし、ここで新しい問いが生まれます。佐藤さんが半年かけて身につけた「人へ任せる方法」を、佐藤さん一人の中だけに残してよいのでしょうか。同じ会社の別の管理職が、また最初から同じ失敗を繰り返し、半年かけて同じことを学ばなければならないのでしょうか。

そして、佐藤さんが経験した失敗も、「本人が学んだからそれでよい」と終わらせてよいのでしょうか。ここから、学習の単位がもう一段上がります。

人が学んだことを、AI と会社の仕組みへ戻す。

次章では、成功から得られた実践知を次の人へ渡す**名人循環(Mastery Loop)**と、失敗から仕組みそのものを改善する**失敗防止循環(Foolproof Loop)**を見ていきます。

第14章 人が育つほど、AIと仕組みも育つ

——名人循環(Mastery Loop)と失敗防止循環(Foolproof Loop)

第8章から、佐藤健太さんの半年間を追ってきました。最初の佐藤さんは、自分で問題を解決し、自分で成果を出せる優秀なプレイヤーでした。しかし、その強さゆえに、重要な仕事ほど自分で判断し、部下へ仕事を任せにくいという課題も抱えていました。

半年間、佐藤さんは週次行動(Weekly Action)を試し、AIとの対話を通じて行動過程(Process)を振り返ってきました。チームでは働き方そのものを見直し、節目には心理状態(Mind)を振り返り、上司は別の経路から成果(Result)を確認しました。そして三つのGapから、本人、上司、チーム、組織のどこを変えるべきかを考えてきました。

その結果、佐藤さんは少しずつ、「自分が成果を出す人」から「人を通じて成果を生み出す人」へ変わり始めました。

しかし、学習する組織(Learning Organization)を目指すなら、ここで終わることはできません。

佐藤さんが半年間かけて身につけたものを、佐藤さん一人の中だけに残してよいのでしょうか。別の管理職が同じ課題に直面したとき、また最初から同じ失敗を繰り返し、同じ時間をかけて学ばなければならないのでしょうか。また、佐藤さんが経験した失敗も、「本人が反省したからそれでよい」と終わらせてよいのでしょうか。

個人が学んだ後には、もう一段上の学習があります。**人が学んだことを、AIと会社の仕組みへ戻すことです。**

本章では、その二つの経路を、成功から学ぶ**名人循環(Mastery Loop)**と、失敗から学ぶ**失敗防止循環(Foolproof Loop)**として整理します。

ただし、これらは私が作った名称で、確立された学術理論の名称ではありません。人とAIと組織が、成功と失敗の両方からどのように学び続けるのかを説明するために、本書で整理している実務モデルです。

AIが「学ぶ」とは、勝手に会話を吸収することではない

最初に、「AIも人から学ぶ」という言葉の意味を明確にしておきます。

本書でいうAIの学習とは、社員がAIとの対話で話した内容が、すべて自動的に取り込まれ、AIが勝手に再学習していくという意味ではありません。

本人との対話には、本人自身の内省のために守るべき情報があります。また、佐藤さんに有効だった方法が、別の人や別のContextでも有効とは限りません。一人の経験をそのまま「会社の正解」にしてしまうことも危険です。

そこで、人の経験から有効そうな知識が見つかって、最初は**知識候補**として扱います。なぜその方法が有効だったのか。どのようなContextで使えるのか。他の人でも再現できるのか。例外はないのか。必要であれば、別の場面でも試します。

そのうえで、人間が「これは組織として利用できる」と判断したものを、AIの質問設計、標準手順、知識データ、業務ルールなどへ反映します。

流れとしては、『**人の経験 → 知識候補 → 確認・試行 → 人間による判断 → AI・標準へ反映**』となります。

AIが組織の正解を自動的に決めるものではありません。AIや標準を更新するときにも、最後の意味づけと判断には人間が入ります。

名人循環(Mastery Loop)——一人が学んだ地点を、次の人の出発点にする

佐藤さんが半年間で学んだことを考えてみましょう。

最初の佐藤さんにとって、「部下へ任せる」という行為は、「全部任せる」か「自分でやる」かに近いものでした。そのため、失敗させられない重要案件ほど、自分で仕事を引き取っていました。

しかし、実際に試行錯誤を重ねる中で、「任せる」という行為にもいくつもの段階があることが分かってきました。最初に本人の案を聞く。なぜそう考えたのか説明してもらう。判断できる範囲を明確にする。最終判断だけ自分が持つ。失敗してもよい範囲を決める。そして仕事の後で一緒に振り返る。

これらは、佐藤さん自身が仕事の中で獲得した実践知です。

しかし会社が、「佐藤さんは良い管理職になった」で終われば、佐藤さん個人は成長しても、会社全体の能力はそれほど上がりません。

そこで次に、佐藤さんが**何を見て、何を問い、どのように判断していたのか**を振り返ります。

例えば、「部下から質問されたとき、すぐ答えを教えず、まず本人の案を聞く」という方法が有効だったとします。

ただし、一度うまくいったからといって、そのまま全社標準にはしません。経験の浅い部下にも有効なのか。緊急時でも同じなのか。法務や安全に重大な影響を与える判断でも使えるのか。他の管理職が使っても機能するのか。その方法を使える条件を確認します。

複数の場面で有効性が確認できれば、AIの問いとして、「本人へ答えを示す前に、本人自身の案を確認できる場面ではありませんか。」という形で組み込むことができます。

すると、新しく管理職になった別の社員は、佐藤さんが何週間も試行錯誤して到達した問いに、もっと早い段階で出会えます。

次の人は、前の人と同じ場所から始める必要がありません。その人はさらに先へ進むことができます。誰に何を任せるのか。どこまで権限を渡すのか。どの判断を標準化できるのか。自分が不在でもチームが判断できる状態をどうつくるのか。そこで新しい実践知が生まれれば、それをまた知識候補として会社へ戻します。

AI や標準の水準が上がる。人間は、その先の仕事へ進む。そこで人間が新しい知識を生み、その知識が AI と標準をさらに引き上げる。

この循環を、本書では**名人循環(Mastery Loop)**と呼びます。一周するたびに、次の人の学習開始地点が少しずつ高くなっていきます。

名人を育てるだけでは、名人依存になる

人材育成には一つの逆説があります。優秀な人が育てば育つほど、その人に仕事が集まることです。

「あの顧客は佐藤さんでなければ対応できない」「あの判断は田中さんしかできない」「あの上司のチームだけ部下が育つ」。こうした状態は、確かに優秀な人がいる会社ではあります。しかし、能力が個人の中に閉じているなら、必ずしも学習する組織(Learning Organization)とは言えません。

学習する組織では、名人を育てるに加えて、名人が持っている知識のうち、他の人へ渡せる部分を組織へ残すことを考えます。

ただし、名人の行動をそのままコピーするわけではありません。熟練した人ほど、相手や状況(Context)に応じて判断を変えているからです。ある部下には問いを返し、別の部下には先に答えを教える。小さな仕事なら任せるが、重大なリスクがある仕事では自分が判断する。

この違いまで含めて熟練です。だから、形式知化するときに本当に取り出したいのは、「この人は何をしたのか」という表面的な行動だけではありません。

「何を見て、その判断をしたのか。」そこまで取り出す必要があります。

そして、何を標準化するのか。何を AI の問いへ変えるのか。何を人間の判断として残すのか。この三つを分けます。

AIの水準が上がることは、人間の判断をなくすことではありません。むしろ、人間が本当に判断すべき場所を明確にすることでもあるのです。

失敗防止循環(Foolproof Loop)——失敗した人を直すだけでは終わらない

人は、成功だけから学ぶわけではありません。失敗には、成功とは別の重要な情報があります。

例えば佐藤さんが、現地社員へ仕事を任せたとします。それまで自分で抱えすぎていた反省から、今回はかなり広い範囲を任せました。

ところが、現地社員が契約上重要な条件を確認しないまま顧客へ回答してしまいました。後から問題が分かり、佐藤さんが修正することになったとします。

佐藤さんは、「任せすぎました。次からもっと注意します」と振り返るかもしれません。これは個人の学習としては意味があります。

しかし、組織の学習としては、それだけでは足りません。もう一つ問いがあります。「なぜ、その間違いは顧客へ届くところまで通過できたのでしょうか。」

ここではまず、故意の違反と意図しないミスに分ける必要があります。

ルールを理解したうえで意図的に破ったのであれば、人の責任として扱わなければなりません。一方、正しく仕事をしようとしていたにもかかわらず起きたミスであれば、本人への注意だけで終わらせません。

確認項目は明確だったのか。誰が最終確認するかという Role は決まっていたのか。経験の浅い社員にも理解できる手順だったのか。システム上、重要項目を未確認のまま処理できてしまったのか。同じ失敗が他でも起きていなかったのか。こうしたことを確認します。

故意の違反は人の責任として扱う。故意でないミスは、仕組みが学ぶ材料として扱う。ここから、失敗防止循環(Foolproof Loop)が始まります。

「次から気をつけます」を、仕組みの変更まで進める

仕事で失敗が起きたとき、「確認不足でした。次から注意します」という再発防止策はよく使われます。

もちろん、本人が注意することにも意味はあります。しかし、人の注意力だけに依存する仕組みは弱いものです。人は疲れます。忘れることもあります。経験の浅い人もいます。同じ条件が残っているなら、別の人が同じ失敗をするかもしれません。

そこで、第1章で見たジェームズ・リーズン(James Reason)のシステムアプローチ(System Approach)が重要になります。

「誰が間違えたのか」だけでなく、「なぜ、その間違いが失敗として通過できたのか」を見るのです。

佐藤さんのケースを調べた結果、契約条件の確認方法が担当者ごとに異なり、重要項目を確認しないまま顧客へ回答できる状態だったとします。

それなら、本人へ注意するだけでなく、仕組みそのものを変える余地があります。

そこで SCAD へ上げます。問題を共有(Share)し、実際に同じ問題が他にもあるのか確認(Check)する。チェック項目や Role、システムを調整(Adjust)し、変更した仕組みを実行(Do)する。そして一定期間後、本当に同じ種類の失敗が減ったのかを再び観測します。

減っていなければ、改善策そのものをもう一度見直します。

つまり、『**失敗 → 構造問題の探索 → SCAD → 仕組み変更 → 再実行 → 再観測 → 再調整**』という循環です。

これが失敗防止循環(Foolproof Loop)です。

人を叱って終わる組織では、失敗は消費されます。**失敗から仕組みを変える組織では、失敗が学習資産になります。**

成功も失敗も、個人の中で終わらせない

名人循環(Mastery Loop)と失敗防止循環(Foolproof Loop)は、一見すると反対の循環です。

名人循環は成功から始まります。優れた実践を知識候補として取り出し、再現できる部分を AI や標準へ戻す。次の人は、前の人より少し高い地点から学習を始めます。

一方、失敗防止循環は失敗から始まります。ミスの背後にある構造を調べ、仕組みを変え、同じ失敗が起きにくい状態をつくります。

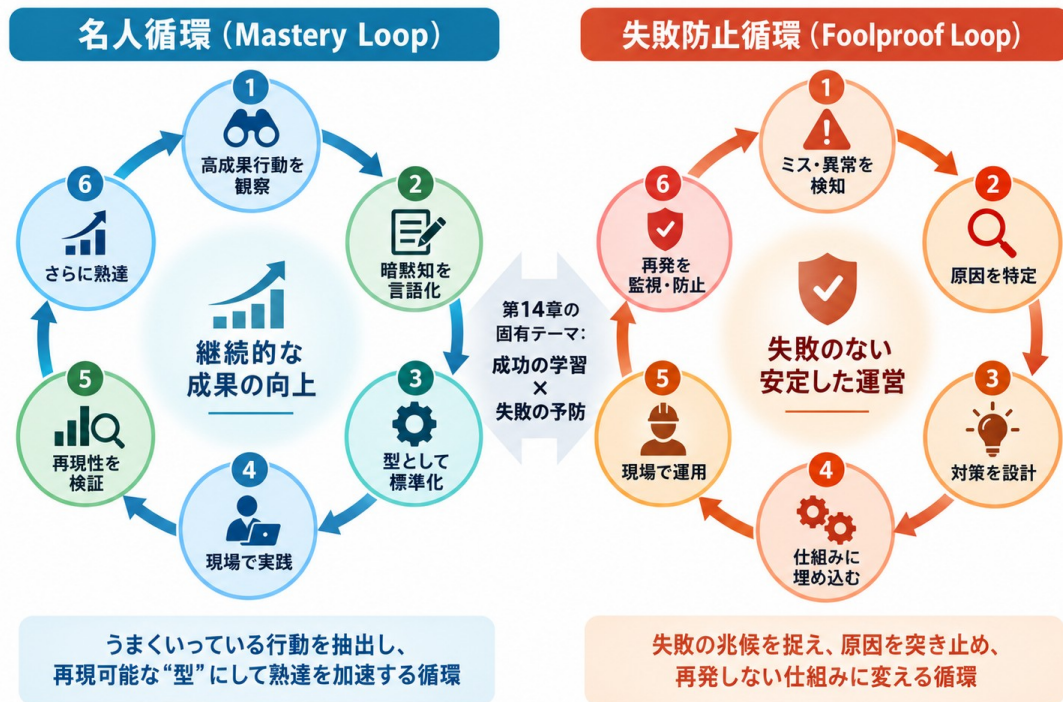
しかし、二つには共通する問いがあります。「**この経験から、会社は何を学べるのか。**」

成功した人がいたら、「優秀な人だった」で終わらせない。何が他の人にも使えるのかを考えます。

失敗した人がいたら、「注意不足だった」で終わらせない。何を仕組みとして変えられるのかを考えます。

成功も失敗も、個人の経験から組織の学習材料へ変える。そこに二つの循環の共通点があります。

図14-1 名人循環 (Mastery Loop) と失敗防止循環 (Foolproof Loop)



SCAD が、個人の学習と組織の学習をつなぐ

ここまで第三部では、SCAD が何度も登場しました。

AI コーチングの中で、本人だけでは変えられない問題が見つかったとき。複数人に同じ Gap が見つかったとき。Role と制度の矛盾が見つかったとき。そして本章では、成功した方法を組織へ広げるときや、失敗を仕組み改善へ変えるときにも SCAD が使われます。

SCAD は、単なる改善提案制度ではありません。**個人の経験を、組織の変更へ渡す経路です。**

EvaliA が個人の気づきを生み出す。その中から組織として扱うべきものを、SCAD が標準、業務、制度、システムの変更へつなぐ。変わった仕組みの中で、また人が仕事をする。そこで新しい成功や失敗が生まれ、再び学習が始まります。

こうして、人材成長循環(Human Development Loop)と組織改善が一つにつながります。

人が仕組みを育て、仕組みが人を育てる

私は長い間、人を育てることに強い関心を持ってきました。人が成長すれば、会社も成長すると考えていたからです。

しかし、人に依存しすぎれば、その人がいなくなった瞬間に能力が失われます。そこで私は、標準化や仕組み化を進めるようになりました。

ところが、仕組みだけを固定すれば、今度は人が仕組みに従うだけになります。現場で新しい知識が生まれても、それが仕組み側へ戻らなければ、組織はやがて硬直します。

人か、仕組みか。そう考えること自体が違っていたのだと思います。

良い仕組みは、人の学習を速くします。成長した人は、現在の仕組みでは扱えない新しい問題に気づきます。その人の経験を使って仕組みを変える。そして、変わった仕組みが、次の人をさらに高い場所から育てる。

人が仕組みを育て、仕組みが人を育てる。

AIが加わることで、この往復を、以前より短い周期で続けられる可能性が生まれました。

半年前の佐藤さんと、今の佐藤さん

第三部の最初に戻ってみましょう。半年前の佐藤さんは、「海外拠点責任者になるなら、自分が一番仕事のできる人でなければならない」と考えていました。

半年後の佐藤さんは、少し違うことを考えています。

「自分が一番多く仕事をする事より、自分がいなくても現地社員が判断し、顧客へ価値を提供できる状態をつくる方が、責任者として大切なものかもしれない。」

本人のものの見方が変わりました。行動も変わりました。チームの仕事の仕方も変わり始めました。

しかし、それだけではありません。佐藤さんが身につけた方法の中から、次の管理職にも使える知識候補が生まれました。失敗した経験からは、会社の仕組みを変えるべき点も見つかりました。

ここで初めて、**佐藤さん一人の学習が、会社の学習へ変わります。**

学習する組織(Learning Organization)は、失敗しない組織ではない

学習する組織(Learning Organization)とは、全員が優秀になり、誰も失敗しなくなる会社ではありません。

人は失敗します。上司も判断を間違えます。KPI の設定が間違っていることもあります。Role の置き方が適切でないこともあります。AI が返した仮説や問いが適切でないこともあるでしょう。

重要なのは、間違いが起きないことではありません。間違いが見えたときに、その経験から学び、次の行動や仕組みを変えられることです。

成功したときも同じです。成功者を表彰して終わるのではなく、なぜ成功したのか、どこまで他の人へ渡せるのかを考えます。

成功から学ぶ。失敗からも学ぶ。そして、その学びを次の人と次の仕組みへ渡す。 それによって、個人の経験が会社の能力へ変わっていきます。

人を育てるのでも、仕組みだけを育てるのでもない

プロローグで、私は自分自身の人材育成を振り返りました。長い間、人を直接育てようとしてきました。その後、人へ依存することの限界から、仕組み化を進めました。

そして今、AIを使った人と組織の学習を考える中で、この二つは一つにつながりました。

人を育てるのでも、仕組みだけを育てるのでもない。人が仕組みから学び、仕組みも人から学ぶ。

成功した経験は名人循環(Mastery Loop)へ戻す。失敗した経験は失敗防止循環(Foolproof Loop)へ戻す。その二つが、AI、SCAD、標準、業務、制度を少しずつ更新します。更新された環境の中で、次の人がまた学び始めます。

しかし、次の循環は、前と同じ場所から始まるわけではありません。**前の人が残した学習の上から始まります。**

ここまで来ると、EvaliAは単なる評価システムでも、AIコーチでもありません。人とチームと会社が、学習を止めないための仕組みになります。

ここまでが設計である

しかし、設計がどれほど整っていても、現場で動かなければ意味はありません。

社員は、本当に毎週AIと対話するのでしょうか。実際の行動は変わるのでしょうか。個人の学習はチームへ広がるのでしょうか。Process、Mind、ResultのGapから、本当に意味のある介入(Intervention)を見つけられるのでしょうか。

そして、もう一つ重要な問いがあります。AIが多くを観測し、多くを問い、多くのパターンを比較できるようになったとき、**人間は何を担うべきなのでしょう。**

AIだけで学習する組織をつくれるのであれば、人間は最後に承認するだけでよいことになります。しかし、本書でここまで見てきた循環は、そうなっていません。

本人が未来を選ぶ。上司が Role を判断する。本人が次の行動を選ぶ。チームが自分たちの働き方を問い直す。上司が成果の意味を判断する。Gap からどこを変えるかを決める。そして、個人の経験を会社の知識や仕組みへ戻すかどうかを人間が判断する。

学習循環の重要な場所には、何度も人間が入っています。

では、なぜ人間が入る必要があるのでしょうか。そして、Human in the Loop は、この学習循環の中でどのような役割を持っているのでしょうか。

次の終章では、第三部で動かしてきた仕組みをもう一度理論へ戻し、**なぜ AI だけでは「学習する組織」はつukれないのか**を考えていきます。

終章 AI だけでは「学習する組織」は つukれない

——Human in the Loop が、人と AI と組織の学習 を循環させる

第三部では、佐藤健太さんという一人の人物を通して、EvaliA が実際にどのように動くのかを見てきました。本人の未来から計画(Plan)をつくり、役割(Role)と重要業績評価指標(Key Performance Indicator)を定め、それを週次行動(Weekly Action)まで小さくする。毎週の人工知能(AI)コーチングでは行動過程(Process)を振り返り、その学びをチームへ広げ、節目では心理状態(Mind)を本人自身が言語化する。さらに上司は別の経路から成果(Result)を確認し、三つの間に生じた Gap から、本人、上司、チーム、組織のどこを変えるべきかを考えました。

そして、成功から得られた実践知は名人循環(Mastery Loop)へ、失敗から得られた教訓は失敗防止循環(Foolproof Loop)へ戻し、一人の経験を会社の知識や仕組みへ変えていくところまで進みました。

ここまでで、EvaliA が「どのように動くのか」はかなり具体的になりました。

しかし、本書を閉じる前に、もう一度最も根本的な問いへ戻る必要があります。

なぜ、これが学習する組織(Learning Organization)なのでしょう。人工知能(AI)が毎週社員へ質問することと、学習する組織とは何が違うのでしょうか。Process、Mind、Result を比較することが、なぜ単なる人事データ分析ではなく、組織学習になるのでしょうか。

その問いを考える中心にあるのが、人間参加型(Human in the Loop)です。

最上位の目的は、EvaliA ではない

ここまで EvaliA について多くのことを書いてきましたが、本書の最終目的は、EvaliA を導入することではありません。

EvaliA は手段です。

最上位にある目的は、**学習する組織(Learning Organization)**をつくることです。そして、その組織が本当に学習する方向へ向かっているのかを見るための定規として、本書ではセンゲの五つのディシプリン(Five Disciplines)を使ってきました。

本人がどこへ向かいたいのかを考える自己マスタリー(Personal Mastery)。自分自身の前提やものの見方を問い直すメンタルモデル(Mental Models)。本人の未来と会社の未来を接続する共有ビジョン(Shared Vision)。個人の経験をチーム自身の学びへ変えるチーム学習(Team Learning)。そして、個人だけでなく、上司、チーム、制度、業務構造まで原因探索を広げるシステム思考(Systems Thinking)です。

EvaliA が、この五つを発明したわけではありません。私たちが試みているのは、これまで経営学の中で示されてきた学習の考え方を、日常の仕事の中で繰り返し運転できる仕組みにすることです。

その具体的な運転構造が、人材成長循環(Human Development Loop)です。EvaliA は、その循環を実装するエンジンです。そして人間参加型(Human in the Loop)は、その上に置かれるもう一つの階層ではありません。**人材成長循環の重要な判断点を貫く運転原則**です。

終章 1

人とAIが共に学び続ける組織へ

— 観測・対話・介入・再観測の継続で、個人・チーム・組織は進化し続ける —



学習する組織を実装する全体構造

Learning Organization を目的に、EvaliA と SCAD で運転する構造整理

階層	名称	意味・役割
最上位目的	学習する組織 (Learning Organization)	個人・チーム・組織が、学び続けながら適応し、成長し続けることを目指す最上位目的
設計を検証する定規	五つのディシプリン (Five Disciplines)	学習する組織として設計が妥当かを点検するための理論的な定規
運転構造	人材成長循環 (Human Development Loop)	個人の学びと成長を、観測・対話・評価・再実行へと循環させる運転構造
実装エンジン	EvaliA	週次観測と対話を通じて、個人の状態・Gap・行動を可視化し、循環を実装するエンジン
組織改善サブグループ	SCAD	個人の課題を、組織構造・役割・標準化・改善へ接続し、組織知へ変える補助ループ

Human in the Loop が貫く判断点

目標設定 ▶ 意味づけ ▶ Role設定 ▶ 行動選択 ▶ 評価 ▶ 介入 ▶ 組織知化

これらの判断点は AI に任せきりではなく、人間参加型（Human in the Loop）が一貫して担う。

— AI は循環を支え、人間は意味づけと判断を担う。 —

Human in the Loop は、「最後に人が承認すること」ではない

Human in the Loop という言葉から、AI が出した判断を、最後に人間が確認する場面を思い浮かべる人もいるでしょう。

もちろん、それも一つの使い方です。しかし、EvaliA で考えている Human in the Loop は、それより広いものです。

人間は最後に一度だけ登場するのではありません。学習循環の途中に何度も入ります。

本人が「何になりたいか」を決める。会社がどこへ向かうのかを決める。上司と本人が現在の Role を確認する。本人が次の Weekly Action を選ぶ。AI が本人の言葉を整理したとき、その意味が合っているかを本人が確認し、必要なら訂正する。上司が Result について判断する。Gap が見つかったときには、本人、上司、チーム、組織のどこへ働きかけるのかを人間が決める。そして、成功した知識のうち何を AI や会社の標準へ戻し、何を人間の判断として残すのかを決めます。

これらはすべて、単なる情報処理ではありません。意味と価値と責任を伴う判断です。

したがって、本書における Human in the Loop とは、**人間が最後に承認する仕組みではなく、学習循環の中で意味づけ、選択、判断、責任を引き受け続ける仕組み**です。

AI が得意なのは、「観測を止めないこと」である

では、AI は何を担うのでしょうか。

AI が会社の理念を決めるわけではありません。社員の未来を決めるわけでもありません。誰を管理職にするのかを決めるものでもありません。

AIが非常に得意なのは、過去を記憶し、現在と比較し、繰り返しているパターンを見つけ、問いを返し続けることです。

一人の社員が三か月前に何を話していたのか。同じところで何週間止まり続けているのか。本人は「変わった」と感じているが、実際の行動も変わっているのか。複数の社員が同じ障害にぶつかっていないか。本人と上司で、同じ出来事の見方が違ってないか。

こうしたことを、多くの社員について、週単位で継続する。優れた上司なら、一人の部下へ良い質問をすることはできるでしょう。しかし、数十人、数百人に対して、過去の対話を踏まえながら継続的に問いを返し続けることは容易ではありません。

AIが大きく変えたのは、「答え」だけではありません。

継続です。

AIは、人間に代わって意味を決めるのではなく、人間と組織の学習が止まらないように、観測、比較、問いを継続する役割を担います。

最適化することと、前提を問い直すことは違う

AIは、与えられた目標に対する Gap を見つけることが得意です。

例えば会社が「売上100」という目標を置き、現在が80だったとします。あと20をどう埋めるのか。どの顧客へ行くのか。どの商品売るのか。どの行動を増やすのか。こうした分析や選択肢生成は、AIが支援しやすい領域です。

しかし、学習する組織(Learning Organization)には、もう一つ重要な問いがあります。

「そもそも、100という目標でよいのか。」という問いです。

目標そのものが間違っているかもしれません。KPIが間違っているかもしれない。Roleが間違っているかもしれない。本人の能力ではなく、商品や制度に問題があるかもしれません。もっと努力するのではなく、その仕事そのものをなくす方が正しい場合もあります。

第1章で取り上げたアージリスの単一循環学習(Single-loop Learning)と二重循環学習(Double-loop Learning)の違いは、ここに 있습니다。

単一循環学習では、基本的な前提を維持したまま行動を修正します。一方、二重循環学習では、「なぜ、その目標なのか」「なぜ、その KPI なのか」「なぜ、その Role なのか」「なぜ、本人を変えなければならないのか」と、前提そのものを問い直します。

Human in the Loop があるだけで、自動的に二重循環学習が起きるわけではありません。人間が AI の提案をそのまま承認するだけなら、前提は変わらないからです。

しかし、AI が Gap を見つけ、それをきっかけに人間が行動だけでなく Role、KPI、制度、さらには目標そのものまで問い直せるなら、二重循環学習(Double-loop Learning)を日常の経営運営へ持ち込むことができます。

五つのディシプリンを、AI と人間でどう回すのか

この視点に立つと、五つのディシプリン(Five Disciplines)と AI、Human in the Loop の関係も整理できます。

表 終-1 五つのディシプリンを、AIと人間でどう回すか

人工知能 (AI) が支援することと、人間が引き受けること

五つのディシプリン	人工知能 (AI) が支援すること	人間が引き受けること
自己マスタリー (Personal Mastery)	過去との比較、現在地やGapを考える問い	何を目指すのかを選ぶ
メンタルモデル (Mental Models)	繰り返すパターンや矛盾を鏡として返す	自分の前提を問い直し、選び直す
共有ビジョン (Shared Vision)	本人の未来と会社の方向の接点を整理する	何を本当に共有するのかを対話して決める
チーム学習 (Team Learning)	共通パターン、Team CAGE、共通Gapを構造化する	異なる視点を使い、Team Actionを選ぶ
システム思考 (Systems Thinking)	個人を越えた共通Gap、反復、構造パターンを示す	原因仮説を検証し、どこへ介入するかを決める

**AIが五つのディシプリンを代行するのではない。
AIが日々の仕事の中で映しやすくし、人間が見る方向を決める。**

AIが自己マスタリーを代行するわけではありません。本人の未来をAIが決めてしまえば、それは本人の自己マスタリーではありません。共有ビジョンも同じです。社員の目標を会社の目標へ機械的に一致させることが共有ビジョンなのではありません。

システム思考についても、AIは「個人だけが原因ではない可能性」を示すことはできます。しかし、どの仮説を採用し、何を変えるのかには人間の判断が必要です。

つまり、AIが五つのディシプリンを代行するものではありません。AIが、五つのディシプリンを日常の仕事の中で回しやすくし、人間がその意味と方向を決めます。

人材成長循環(Human Development Loop)に、人間参加型を組み込む

第三部で見てきた人材成長循環(Human Development Loop)を、この観点からもう一度整理します。

計画(Plan)では、本人がやりたい姿を考え、会社が方向を示し、上司と本人が Role を確認して KPI を置きます。実行(Do)では、本人が Weekly Action を選び、実際の仕事の中で試します。

確認(Check)では、AI が Process を継続的に構造化し、節目の Mind、上司側の Result という別の経路と比較できるようにします。改善(Act)では、AI が原因を確定するのではなく、Gap を材料に、人間が本人、上司、チーム、組織のどこへ介入するのかを判断します。必要であれば SCAD へ送り、制度、標準、業務、システムを変更し、その結果をもう一度観測します。

つまり、人材成長循環(Human Development Loop)とは、単に PDCA を人材育成へ置き換えたものではありません。**PDCA のそれぞれの判断点に Human in the Loop を組み込み、個人、チーム、組織が前提まで含めて学び直せるようにする循環**だと、私は考えています。

本書の設計では、Human in the Loop は「学習を成立させる運転原則」である

AI について Human in the Loop が語られるとき、「AI は間違えるので、最後は人間が確認しましょう」という安全対策として説明されることがあります。

もちろん、その役割も重要です。しかし、本書で考えている Human in the Loop の意味は、それだけではありません。

人間が入らなければ、何を目指すのか、何を成果と考えるのか、どの Gap に意味があるのか、どの価値を優先するのか、本人を変えるべきなのか、それとも会社を変えるべきなのか、といった問いを引き受ける主体がいなくなります。

これらは、単なる情報処理ではありません。「意味と責任」を伴う判断です。

したがって、本書の設計では、**Human in the Loop** は、**AI利用の最後に置く安全装置だけではなく、人間が意味づけ、選択、判断、責任を担いながら、学習循環を成立させる運転原則**として位置づけます。

では、本当にこの循環は動き始めたのか

ここまでは理論と設計です。しかし、理論がどれほど整っていても、実際に社員が使わず、行動も組織も変わらないのであれば意味がありません。

そこで最後に、現在までの初期実証を見ていきます。

ただし、以下は完成した効果検証ではありません。当社で **2026 年 9 月時点で運用中の EvaliA から得られた初期実証のスナップショット**です。

また、異なる目的で取得されたデータを混同しないように分けて扱います。第一は、2026 年 7 月 5 日の週から 9 月 13 日の週までの 11 週間、325 名について蓄積された 1,906 件の週次所見です。第二は、2026 年 8 月 3 日から 8 月 24 日までの 4 週間、4 週連続で参加した 149 名を追った時系列補助分析です。第三は、日本の 58 名について、第 4 回時点の AI との対話の使われ方を分析した対話品質補助分析です。さらに、上司介入については、2026 年 8 月 16 日から 9 月 13 日までの上司 14 名、対象者 56 名、80 件の記録を補助的に使っています。

実証スナップショット【更新基準日：2026 年 9 月 13 日】

表 終-2

EvaliA 11週間の運用スナップショット
更新基準日：2026年9月13日

週次観測・進捗分類		原因主体分類 (n=1,906)	
主分析対象期間	11週	原因主体：本人	385件
対象者	325名	原因主体：本人+上司	319件
週次所見	1,906件	原因主体：業務構造	83件
前進	1,187件	原因主体：会社・部門	26件
停滞	716件	原因主体：上司	17件
後退	3件	原因主体：未判定	1,076件
明瞭度が閾値未満	627件 (33%)	上司記録	
本人の自己説明を含む所見	737件	上司記録	80件
		上司記録対象	上司14名・対象者56名
		上司対応済み	61件
		上司未対応	10件
		上司未入力	9件

この数値が示すのは「人が何%成長したか」ではない。
「前進・停滞・後退」は週次所見の分類であり、人格や能力全体の成長を意味しない。
確認できるのは、11週間にわたり観測・分類・上司記録が継続的に蓄積されたという運用実績である。

※進捗判定基準は2026年8月27日に変更されているため、本集計は同一基準による時系列的な効果比較を示すものではなく、当該期間における運用スナップショットである。
※一部指標は週次所見1,906件を母数としている。

ここで注意したいのは、「前進 1,187 件だから、約 6 割の社員が成長した」とは言えないことです。1,187 は人の数ではなく週次所見の件数です。また、「前進」はシステム上で観測された週次の行動状態であり、その人の人格や能力全体が成長したことを意味しません。

この数値から確認できるのは、一回限りのサーベイではなく、人の行動状態を週単位で継続して観測する仕組みが、実際の運用の中で動いているということまでです。

「分からない」と残せる AI であること

実証データの中で、私が重要だと考えている数字の一つが、明瞭度が閾値を下回った 627 件です。全体の 33%でした。

ここでいう明瞭度は、「AI が正解している確率」ではありません。対話から、どの程度明瞭に行動シグナルを読み取れたのかを見る内部指標です。本分析では、システム内で設定した基準値を下回るものを「閾値未満」としています。

人間の対話は、いつも明確ではありません。回答が短いこともあります。理由が分からないこともあります。Contextが足りないこともあります。本人自身が、まだ考えを整理できていないこともあります。

そのときにAIが無理に、「この人は停滞しています」「原因は上司です」と結論を出す必要はありません。

追加の問いをする。次週も観測する。人間が確認する。あるいは、「現時点では判断できない」と残す。

AIが、分からないものまで分かったように扱わないことも、Human in the Loopにとって重要です。

週次の状態は、一直線には動かない

時系列補助分析では、2026年8月3日から8月24日まで、4週連続で参加した149名について、週ごとの状態変化を追っています。

61名は4週連続で「前進」でした。一方、15名は「停滞→前進→前進→前進」、13名は「前進→前進→停滞→前進」、10名は「前進→停滞→停滞→前進」、さらに10名は「前進→停滞→前進→前進」という動きを示しました。

4週連続参加者の状態変化パターン（2026年8月3日～8月24日）

時系列補助分析では、4週連続で参加した149名について、週ごとの状態変化を追いました。

No.	状態変化パターン	第1週 8月3日	第2週 8月10日	第3週 8月17日	第4週 8月24日	人数 (名)		割合 (%)
1	前進 → 前進 → 前進 → 前進 (4週連続で前進)	前進 →	前進 →	前進 →	前進	61		40.9
2	停滞 → 前進 → 前進 → 前進	停滞 →	前進 →	前進 →	前進	15		10.1
3	前進 → 前進 → 停滞 → 前進	前進 →	前進 →	停滞 →	前進	13		8.7
4	前進 → 停滞 → 停滞 → 前進	前進 →	停滞 →	停滞 →	前進	10		6.7
5	前進 → 停滞 → 前進 → 前進	前進 →	停滞 →	前進 →	前進	10		6.7
6	その他のパターン	—	—	—	—	40		26.8
合 計		—	—	—	—	149		100.0

（注）割合は、4週連続参加した 149名に対する構成比です。

4週連続で「前進」した人は61名（40.9%）でした。

一方で、一時的に「停滞」を経験しながらも、その後「前進」に転じた人も一定数見られました。

このデータから、「人間の成長とはこういうものである」と一般化することはできません。しかし少なくとも、週次で観測される状態は、いつも一直線には動いていません。

前へ進んだ次の週に止まることもある。停滞した後に動くこともある。一度動いた後に、再び止まることもあります。

だからこそ、一回の評価だけで人を決めるのではなく、時間軸の中で見る必要があります。

人はAIから「答え」をもらっているだけなのか

次に、AI との対話そのものを見ます。

日本の 58 名について、第 4 回の対話を再判定した分析では、「問いを受けて考えを進める」が 27 名、「自分の仮説を持ち込む」が 17 名、「最小限の応答」が 12 名でした。また、対話の型では 58 名中 36 名が、最初から最後まで本人が話す「持続」型と判定されています。なお、この三分類の合計は 56 名であり、残る 2 名はこの三分類には含まれていません。

第4回対話の再判定結果(日本・58名)

日本の58名について、第4回の対話を再判定した分析結果

1. 問いに対する応答の進み方

区分	人数(名)	割合(%)	
問いを受けて考えを進める	27	46.6	
自分の仮説を持ち込む	17	29.3	
最小限の応答	12	20.7	
計	56	96.6	

※割合は58名を母数として算出。上記3分類の合計は56名のため、合計は100%にならない。

2. 対話の型

区分	人数(名)	割合(%)	説明
持続型	36	62.1	最初から最後まで本人が話す型
その他の型	22	37.9	持続型以外
計	58	100.0	—

58名中36名(62.1%)が『持続型』と判定された。
また、応答の進み方では『問いを受けて考えを進める』が27名で最も多かった。

この結果だけで、「深い学習が起きた」と証明することはできません。

しかし少なくとも、AIが答えを出し、社員がそれを受け取るだけの使い方しかされていないわけではありません。問いを受けて自分の考えを進める人がいる。最初から自分の仮説を持ち込み、AIを壁打ち相手として使っている人もいます。

AIが考える主体となり、人間が受け身になるのではなく、**人間が考えるためにAIを使っている例が、実際の対話の中に現れ始めている**ということです。

自己認識と実際の対話を、最初から混ぜない

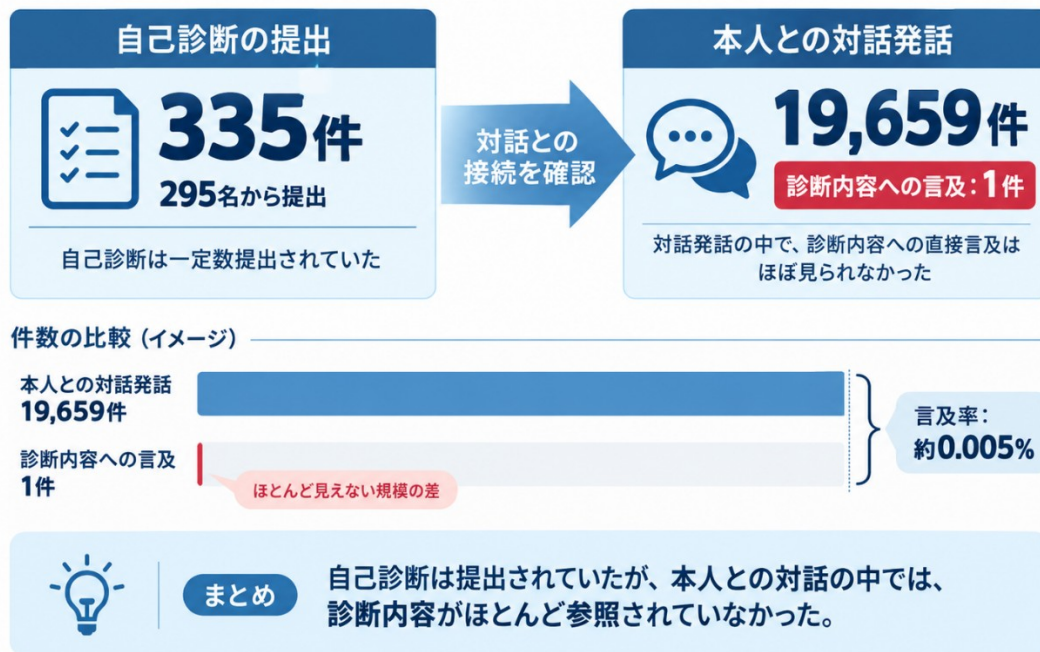
EvaliAでは、一つの診断結果から本人のすべてを説明しないことを重視しています。

実測では、自己診断は335件、295名から提出されていました。一方、本人との対話発話19,659件のうち、診断内容への言及は1件でした。

なお、ここでの自己診断とは、CAGE 適応モデル(CAGE Adaptation Model)の四つの行動機能について、本人が自分の傾向を回答した自己診断を指しています。

自己診断提出と対話内参照のギャップ

実測データにみる、提出量と対話内言及の大きな差



この数字だけで、「自己認識と行動データが統計的に独立している」と証明できるわけではありません。

しかし少なくとも、本人が毎回診断結果を持ち込み、「私はこのタイプだから、このように行動する」という形で対話が自己強化されている状態ではありませんでした。

最初に貼られたラベルが、その後の対話をすべて支配しないようにする。そのためにも、自己認識と実際の行動は、できるだけ別の経路から取り、後から比較する。

これは、第三部で Process、Mind、Result を別々に取り、後から Gap を見る構造と同じ考え方です。

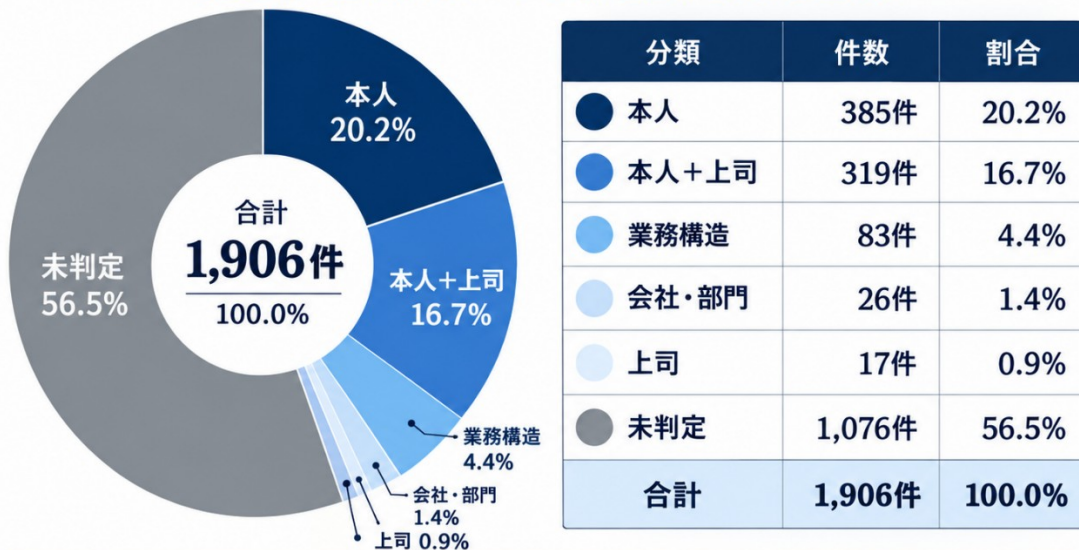
「本人を変える」以外の介入先は、本当に見えてくるのか

学習する組織にするためには、すべての問題を本人の努力へ戻してはいけません。

1,906 件の週次所見について、原因主体は、本人 385 件、本人＋上司 319 件、業務構造 83 件、会社・部門 26 件、上司 17 件に分類され、1,076 件は未判定でした。

週次所見1,906件の原因主体分類

1,906件の週次所見について、原因主体を分類した結果



最多は「未判定」1,076件 (56.5%) であった。一方、判定できた所見では「本人」385件 (20.2%) と「本人＋上司」319件 (16.7%) が中心であった。

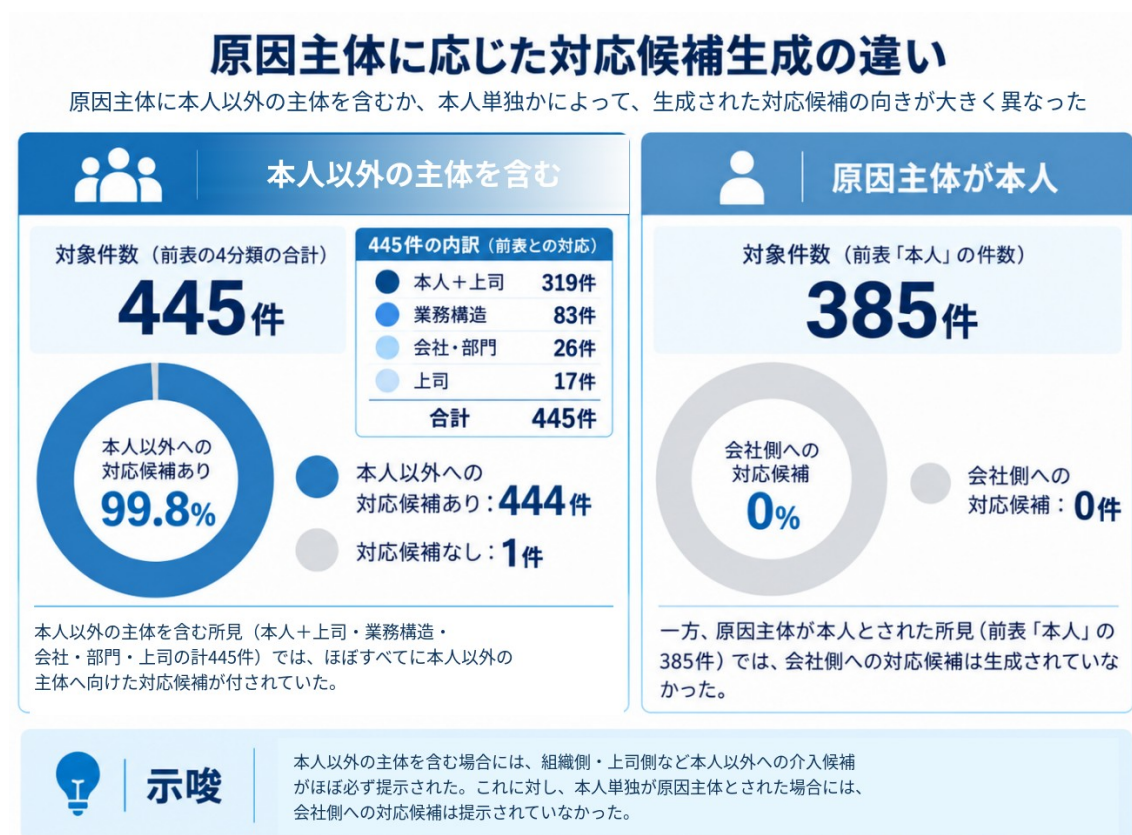
ここには二つの意味があります。

一つは、原因探索が本人の外へ広がっていることです。本人と上司の関係、業務構造、会社や部門、上司自身など、本人以外の可能性も探索対象になっています。

もう一つは、1,076 件が未判定のまま残っていることです。

これは現時点での限界です。しかし同時に、「分からないものまで本人の責任として処理しない」ことでもあります。

さらに、原因主体に本人以外の主体を含む 445 件（本人＋上司 319 件、業務構造 83 件、会社・部門 26 件、上司 17 件）のうち、444 件には本人以外の主体へ向けた対応候補が付されていました。一方、本人が主体とされた 385 件では、会社側への対応候補は生成されていませんでした。



これは、その対応が実際に有効だったことを証明する数字ではありません。しかし、**問題の所在によって、次の介入先を変える AI のロジックが実際に動いていることは確認できます。**

本人で変えられるなら Weekly Action へ返す。本人以外に問題があるなら、上司、チーム、業務、組織側へ視点を移す。

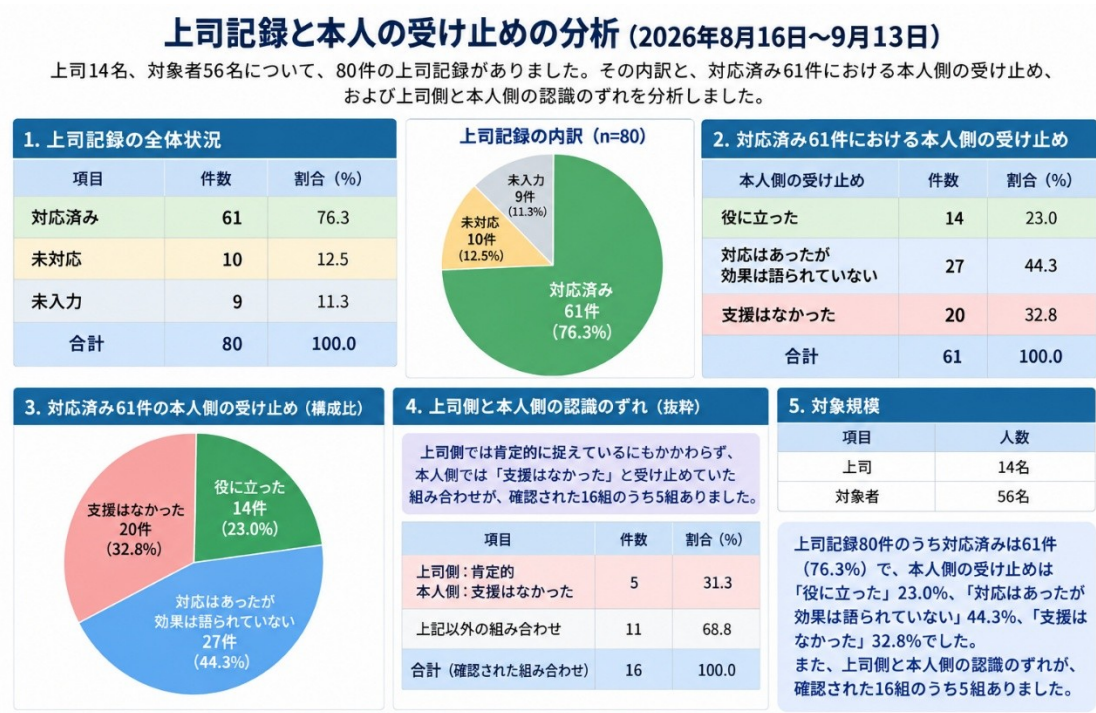
人材育成を「すべて本人を教育すること」にしないための基本構造です。

上司が「支援した」と思っている、本人には届いていない

Human in the Loop の必要性をさらに示すのが、上司と本人の認識差です。

2026 年 8 月 16 日から 9 月 13 日までに、上司 14 名、対象者 56 名について 80 件の上司記録があり、その内訳は対応済み 61 件、未対応 10 件、未入力 9 件でした。

対応済み 61 件について本人側の受け止めを見ると、「役に立った」が 14 件、「対応はあったが効果は語られていない」が 27 件、「支援はなかった」が 20 件でした。さらに、上司側では肯定的に捉えているにもかかわらず、本人側では「支援はなかった」と受け止めていた組み合わせが、確認された 16 組のうち 5 組ありました。



上司は、「自分は部下を支援した」と考えている。しかし本人は、「支援はなかった」と感じている。

どちらか一方だけを真実とする必要はありません。

見るべきなのは、**なぜ二人の認識が違うのかという Gap**です。

ここで重要な機能は、AIを使わないマネジメントでは、上司の「対応済み」で終わってしまうことです。実際は、上司は対応したと思っているのに、部下の立場では、「対応してもらっていない、支援はなかった」と思っています。この事実が表面化することは、ほとんどありませんでした。AIをマネジメントの中に実装したことで初めて分かった事実です。

上司から見れば十分に権限を渡したつもりでも、本人には何を決めてよいのか伝わっていないのかもしれませんが。本人は支援を期待していたが、上司は自律を促すために、あえて介入しなかったのかもしれませんが。

違いそのものが、次の対話の材料になります。

人を観測する AI の「ものさし」も、学習対象になる

実証で学ぶのは、社員や上司だけではありません。AIの判定基準そのものについても、問い直す必要があります。

ある一週間の実測では、毎週仕事内容が変化しやすい層で「前進」と判定された割合が88%だった一方、ルーティン業務を中心とする層では45%でした。

これは一週分の観測であり、職種構成などの影響もあるため、この差だけで「AIに職種バイアスがある」と結論づけることはできません。

しかし、同じ判定基準が職種特性によって異なる結果を生む可能性を示すシグナルではあります。

もしAIが、「新しいことを試した＝前進」と強く見るのであれば、毎週仕事内容が変化する職種ほど前進と判定されやすくなり、安定して正確に繰り返すことに価値がある仕事は不利に見える可能性があります。

そこで進捗判定基準そのものも見直され、2026年8月27日には「前進」の判定基準が改定されています。ここは重要です。

AIは、人を観測する側です。しかし、そのAIの「ものさし」自体も、観測され、問い直されなければならない。当社では、これを継続的に行い、AI自体のプログラムを継続的に状況（Context）に合わせて修正しています。

Human in the Loop は、人間を AI から守るためだけではありません。AI の基準そのものを学習させるためにも必要なのです。

現時点で確認できたことと、まだ言えないこと

2026 年 9 月時点の初期実証から、少なくとも、AI との対話を一回限りではなく複数週の時系列情報として蓄積できていること、問いを受けて本人が考えを進めたり自分の仮説を持ち込んだりする使い方が観測されていること、原因探索が本人だけではなく上司、業務構造、会社・部門などへ広がり始めていること、上司と本人の認識が一致しないケースが実際に存在すること、そして AI 側の判定基準そのものも改善対象になることは確認されています。

一方で、現時点のデータから言えないことも多くあります。

EvaliA を使ったことによって会社の業績が何%向上した、という因果効果はまだ示せません。離職率、生産性、利益率がどの程度改善するのかについても、長期データが必要です。対照群を置いた無作為な実験でもありません。観測期間中には進捗判定基準そのものも変更されています。原因主体が未判定の所見も多く残っています。

また、心理状態(Mind)と成果(Result)を含めた三経路すべてについて、長期間、同じ母集団でどのような Gap が生まれ、それに対する介入が最終成果へどう影響したのかという検証も、これから蓄積する必要があります。

CAGE 適応モデル(CAGE Adaptation Model)についても、現在は実務モデルとして利用していますが、測定モデルとしての妥当性や信頼性は今後の研究課題です。

だから、現時点で、「EvaliA によって学習する組織が実証された」とは言いません。

現在確認できているのは、**人と組織の学習循環を継続的に観測し、異なる視点の Gap を見つけ、本人だけではなく上司、チーム、組織へ介入先を変える仕組みが、現場で動き始めているという段階です。**

実証は、EvaliA が正しいことを証明するためだけに あるのではない

実証というと、自分たちがつくった仕組みの正しさを証明するために、データを集めるように思うかもしれません。

しかし、それでは本書で考えてきた学習する組織(Learning Organization)と矛盾します。

実際に動かした結果、AI が十分な明瞭度で判断できないケースがあることが分かった。原因主体を決められないケースも多く残った。上司と本人の認識が一致しないことも分かった。職種によって進捗判定の出方が違う可能性も見つかった。

それなら、EvaliA 自身を変えなければなりません。判定基準を変える。問いを変える。情報の取り方を変える。人間が確認すべき範囲を変える。

つまり、EvaliA は学習する組織をつくるための完成された答えではありません。**EvaliA そのものも、人と組織から学び続ける必要があります。**

第 14 章で説明した名人循環(Mastery Loop)と失敗防止循環(Foolproof Loop)は、社員だけに適用されるものではないのです。

人間が目的を持ち、AI が学習を止めない

本書の出発点は、「人をどう育てるか」という問いでした。

私は長い間、人を直接変えようとしてきました。その後、人に依存することの限界から、標準化や仕組み化を進めました。そして AI が登場しました。

AI を使えば、もっと効率よく人を育てられる。そう考えることもできます。

しかし、ここまで考えてきて、私は少し違う結論にたどり着きました。

AIの役割は、人間の代わりに「正しい人材」をつくることではありません。人間の未来を決めることでもありません。人間を採点することでもありません。

人間が目的を持つ。AIが問いを返す。人間が試す。AIが変化を観測する。人間が意味を考える。AIが別の視点を返す。人間が何を変えるかを決める。そして、その人間の成功と失敗から、AIと会社の仕組みもまた変わっていく。

この往復こそが重要なのです。

AIが変えたのは、「答え」ではなく「運転可能性」

AIだけでは、学習する組織はつくれません。何をを目指すのかを問い直し、自分のものの見方や、会社と自分の関係、チームの働き方、組織の前提そのものを問い直す。その意味と責任を引き受けるのは人間です。

一方で、人間だけでは継続が難しかったことがあります。全社員へ問いを返し続けること。数週間、数か月前の発言と現在を比較すること。複数人に共通するGapを見つけること。一人ひとりの学びを、チームや組織へつなぎ続けることです。

AIによって、これらを以前より低い運転コストで継続できる可能性が生まれました。AIが学習する組織をつくるではありません。人間が学習循環をつくり、AIがその継続を支えるのです。

エピローグ 人を変えることから、人と会社が学ぶことへ

——「AIで簡単につくれた！」の本当の意味

二十年前の「二つのベクトル」

二十年以上前、私は一枚の図を描きました。

一つは、社員本人が向きたい方向です。もう一つは、会社が向きたい方向です。この二つのベクトルが重なるほど、個人の力と会社の力を同じ方向へ向けることができる。私はそう考えていました。

その考え自体は、今でも間違っていないと思います。

ただし、当時の私は、その二つをどう重ねるかについて、今とは少し違う考えを持っていました。会社が目指しているものを社員へ伝え、会社が求める考え方を理解してもらい、その方向へ本人にも成長してもらう。そのために、経営者や上司は社員を教育する必要があると考えていました。

それは、私自身の経験とも重なっていました。私は若いころ、厳しく教えられの中で多くのことを学びました。それまで正しいと思っていた価値観を疑うことも、自分とは違う考え方を受け入れることも、そこで覚えました。だから、自分が教育によって変わることができたのなら、人も教育によって変わることができる考えたのです。

実際、大きく変わった社員もいました。入社時には経験がほとんどなかった若者が、数年後には顧客から頼られる存在になる。自分の仕事だけを見ていた人が、部下やチームの成長まで考えるようになる。そういう姿を何度も見てきました。

その「成功体験」が、私の考えを強くしました。

しかし今振り返ると、そこには大きな見落としがありました。この成功体験こそが、次のステージで「サクセス・トラップ」なるのです。

本人の成長を願うことと、本人をこちらが考える方向へ変えようとすることは、同じではありません。

目的が善意であれば、方法まで正しいとは限りません。私は、このことに気づくまで長い時間がかかりました。

人を変えるのではなく、学べる条件をつくる

以前の私は、社員が期待どおりに動かなかったとき、「どうすれば、この人を変えられるのだろう」と考えていました。

今なら、少し違う問いを持ちます。

この人は、本当は何になりたいのだろう。会社は、この人に何を期待しているのだろう。本人は何を試し、どこで止まっているのだろう。止まっている原因は、本当に本人だけにあるのだろうか。上司の関わり方、Role、評価の仕方、仕事の進め方、会社の制度を変えた方がよいのではないだろうか。

人の成長を諦めたわけではありません。むしろ、その反対です。

人を直接変えようとするをやめたことで、人が自分自身で変わっていく過程を、以前よりよく見られるようになりました。

人は、誰かから正解を教えられたときだけに変わるわけではありません。自分で選ぶ。実際にやってみる。うまくいかない。なぜなのかを考える。別のやり方を試す。その途中で、「自分はずっと、こう考えていたのか」と、自分自身のものの見方に気づく。そして、もう一度選ぶ。

会社がつくるべきなのは、社員を一つの理想像へ近づける教育だけではありません。本人が考え、試し、失敗し、振り返り、また選び直せる条件です。

そして、本人一人では変えられない問題があるのなら、会社の方が変わらなければなりません。

ここで、人材育成と組織づくりは別のものではなくなりました。

人が会社から学ぶ。そして、会社も人から学ぶ。では、AI は、この関係をどう変えたのでしょうか。

AI が変えたのは、人材育成の理論ではなかった

自己マスタリー(Personal Mastery)も、メンタルモデル(Mental Models)も、共有ビジョン(Shared Vision)も、チーム学習(Team Learning)も、システム思考(Systems Thinking)も、AI が発明したものではありません。

PDCA も、目標管理も、研修も、1on1 も、コーチングも、フィードバックも、エンゲージメント・サーベイも、人事評価も以前からありました。

問題は、それらが**バラバラに行われていた**ことです。

研修を受ける。しばらくして上司との 1on1 がある。別の時期にサーベイへ答える。半期に一度、評価を受ける。それぞれには意味があります。

しかし、研修で本人が何に気づいたのか。その後の仕事で何を試したのか。なぜできなかったのか。上司はどのような支援をしたのか。本人は現在の仕事をどう受け止めているのか。最終的に会社が期待した成果は出たのか。

それらは、一つの時間軸として十分にはつながっていませんでした。そして、つなげようと思っても、最大の制約がありました。

人間には時間がないのです。

一人の優れた上司なら、一人の部下へ良い質問をすることはできます。しかし、何十人、何百人という社員に対して毎週問いかけ、前回何を話したのかを覚え、数週間前の自分と現在の自分を比較し、一人ひとりの状態に応じて次の問いを変え続けることは、現実には困難でした。

さらに、一人ひとりの学習を並べ、「同じ場所で何人も止まっていないか」「これは個人の問題ではなく、会社の仕組みの問題ではないか」と見ることまで、人間だけで継続することは難しかった。

AIが変えたのは、理論ではありません。「運転可能性」、つまり運用の方法です。

これまで「やった方がよい」と分かっているながら、人間の時間と記憶の限界によって続けられなかったことを、日常の仕事の中で回し続けられるようになった。

私は、ここにAIの最も大きな意味があると考えています。

バラバラだった人材育成が、一つのPDCAになった

AIによって大きく変わったのは、研修、上司からのコーチング、本人の振り返り、サーベイ、フィードバック、目標管理、成果評価を、**一つの人材成長循環(Human Development Loop)**として接続できるようになったことです。

計画(Plan)では、本人が何を目標しているのかを明確にします。それを会社が目指す方向とつなぎ、RoleとKPIを置きます。

実行(Do)では、Roleを一週間で試せる Weekly Action まで小さくし、本人が実際の仕事の中で試します。

確認(Check)では、AI週次コーチングによって行動過程(Process)を継続して観測します。一定の節目では、本人が仕事をどう受け止めているのかを心理状態(Mind)から見ます。そして上司には別の経路から問いを返し、RoleとKPIに照らして成果(Result)を確認します。

改善(Act)では、三つを後から比較し、そのGapを見ます。問題が本人にあるのか。上司の関わり方にあるのか。チームのRoleや協働関係にあるのか。それとも、会社の制度や業務そのものにあるのか。

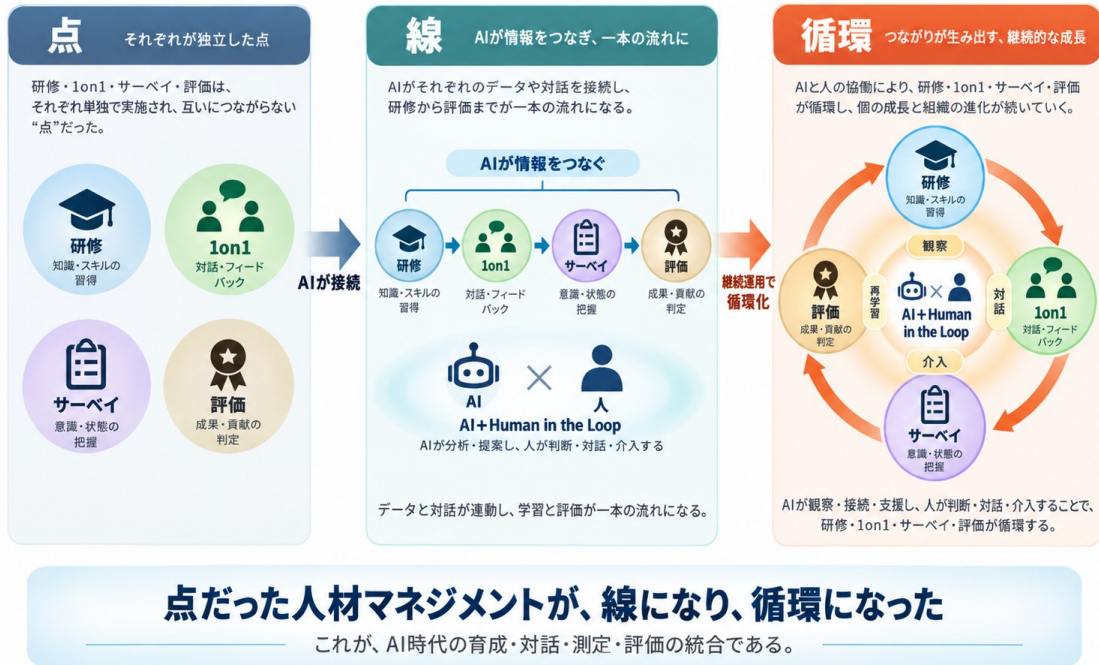
そして、その結果を次のPlanへ戻します。

これまで別々の時期に、別々の担当者が、別々の目的で行っていた人材育成が、**Human in the Loop**によって一つのPDCAとして回り始めた。

これが、本書で考えてきた EvaliA の本質です。

エピローグ1 点→線→循環 AI+Human in the Loopが人と組織をつなぐ

バラバラだった取り組みが、つながり、まわりはじめる。AIと人の協働で、個の成長と組織の進化を持続的に実現する。



AI 週次コーチングが、「続けられない」を変えた

この循環の中でも、特にAIの意味が大きいのが週次コーチングです。

上司が部下へ毎週問いかける。その人が前週に何をすると決めたのかを覚えている。実際にできたのかを聞く。できなかったなら、なぜできなかったのかを聞く。うまくいったなら、なぜうまくいったのかを聞く。そして、次の一週間で何を試すのかを本人に考えてもらう。

優れたコーチならできます。

しかし、全社員について、毎週続けることはできませんでした。

AIによって、この制約が大きく変わりました。

前週の対話を踏まえて問いを返す。三か月前と現在を比較する。同じところで何週間止まっているかを見る。本人自身も気づいていない変化を鏡として返す。

AIが大きく変えたものは、「答え」だけではなく「継続」です。

AIが優れた上司の代わりになったのではありません。人間だけでは物理的に続けられなかった問いと振り返りを、全社員に対して継続できるようにしたのです。

目標管理が、「会社の目標を割り当てること」から変わった

目標管理にも、大きな変化があります。従来の目標管理では、会社の目標を部門へ下ろし、その目標をさらに個人へ割り当てる方法が中心になりやすいでしょう。

もちろん、会社には必要な目標があります。しかし、それだけでは本人にとって、目標は「会社から与えられたもの」になりがちです。

EvaliA では、順番を変えました。

最初に聞くのは、「あなたは、何になりたいのですか。」です。

本人自身も、最初から明確な答えを持っているとは限りません。そこで AI が問いを返します。なぜそうなりたいのか。その状態では何ができるようになっていたいのか。現在の自分との間に何があるのか。

そして初めて、会社の未来との接点を探します。

会社の目標へ本人を合わせるのではありません。本人の未来と会社の未来がどこで重なるのかを考え、その接点を Role へ変えます。

二十年前に描いた二つのベクトルを、AI との対話によって継続的に見直せるようになった。ここにも、AI を使う意味があります。

サーベイが、「点数を聞く」ものから「原因を探す」ものへ変わった

社員の状態を見るサーベイにも、大きな変化が生まれました。

従来のサーベイでは、「満足していますか」「会社を勧めたいですか」といった質問へ、5段階や10段階で回答してもらう方法が中心です。それによって部署間や時期間の比較はできます。

しかし、満足度が3だったとしても、**なぜ3なのか**は分かりません。

AIとの対話では、本人の定性的な回答をさらに掘り下げることができます。

何に満足しているのか。何に不満を持っているのか。何に仕事の意味を感じているのか。何を負担に感じているのか。何を成長と感じているのか。

AIが問いを重ね、その定性的な回答を構造化することで、動機づけ要因、衛生要因、エンゲージメントに関わる状態を、複数の観測指標として時系列で見ることができるようになります。

重要なのは、これらを一つの総合心理スコアへまとめることではありません。点数から本人を推測するのでもありません。

本人の言葉から始め、その意味を確認しながら構造化し、変化を比較できる形にする。

よって、単に「満足度が低い」で終わらず、「本人が不満なのは仕事内容ではなく評価基準が分からないからではないか」「負担感の原因は本人の能力ではなく、Roleが一人へ集中しているからではないか」と、その背後にある本質的な原因仮説まで進むことができます。

ただし、AIが原因を断定するわけではありません。そこから人間が事実を確認します。サーベイが、社員の状態を測るだけのものから、**次にどこを変えるべきかを探す仕組み**へ変わったのです。

CAGEによって、「個人」ではなく「チームの頭」 を見る

CAGE 適応モデル(CAGE Adaptation Model)によって、もう一つ見えるようになったものがあります。

チームです。

重要なのは、「この人は Challenger 型」「この人は Guardian 型」と人を分類することではありません。

成果(Challenger)、探究(Adventurer)、安全(Guardian)、関係(Empathizer)という四つの行動機能が、現在のチームの中で実際に使えているのかをみることです。

成果へ進む力は強い。しかし安全確認が弱い。新しいアイデアは出るが、最後まで実行する行動が弱い。関係は良いが、必要な異論が出ない。

こうしたチームの状態を Team CAGE Map として可視化し、Role や Team Action を見直していく。

私は、これを、四つの異なる機能を状況に応じて使える、いわばホールブレイン的（ハーマンモデル）な状態をチームとしてつくる仕組みだと考えました。

個人に多様性があるだけでは足りません。その違いを、チームの判断能力として使えることが重要です。

評価する上司自身も、学習するようになった

AI によって変わったのは、評価される社員だけではありません。

評価する上司もまた、学習するようになります。

上司にもメンタルモデルがあります。「この社員は優秀だ」「この社員は主体性がない」「管理職とはこうあるべきだ」という過去の印象や前提を持っています。

そこで AI は、社員ではなく上司へ問いを返します。

その評価をした具体的な事実は何か。反対の事実はなかったか。本人がコントロールできない要因はなかったか。上司自身の関わり方が、本人の成果を阻害していなかったか。その成果は再現できるものなのか。

AI が評価を決めるのではありません。

評価する人間自身が、自分の評価の仕方を問い直す。

EvaliA で育つのは、部下だけではありません。上司もまた、学習するのです。

一人の問題から、会社の構造が見えるようになった

さらに AI を使うことで、一人ひとりを深く見ることと、組織全体を見ることを同時に行いやすくなりました。

一人の社員が「承認が遅い」と言っているだけなら、その人固有の問題かもしれません。

しかし、同じ部署の多くの人が同じ場所で止まっているなら、意味は変わります。

一人の中では「本人の主体性」の問題に見えていたものが、複数人を並べることで、「権限が曖昧なのではないか」「承認制度が遅いのではないか」「上司が判断を取りすぎているのではないか」という共通 Gap へ変わります。

本人で変えられるなら Weekly Action へ戻す。チームで変えられるなら Team Action へ。組織の仕組みに問題があるなら SCAD へ上げる。

AI によって、人を細かく見ることと、構造を大きく見ることがつながりました。

成功も失敗も、個人の中で終わらなくなった

そして、学習は評価で終わりません。

ある社員が優れた方法を見つけたなら、「優秀な人だった」で終わらせない。何を見て、何を問い、どのような条件でその判断をしたのかを知識候補として取り出します。他の人にも使えると確認できたものは、AIの問い、標準、業務ルールへ戻します。

これが名人循環(Mastery Loop)です。

一方、失敗が起きたら、「次から注意します」で終わらせない。なぜそのミスが結果まで通過できたのかを問い、必要であればRole、手順、標準、システムを変える。

これが失敗防止循環(Foolproof Loop)です。

人が学ぶ。会社がその人から学ぶ。変わった会社の中で、次の人が前の人より高い場所から学び始める。**人が仕組みから学び、仕組みも人から学ぶ。**

ここまで来て初めて、個人の成長が組織の能力へ変わります。

AIを使わない場合と、究極的に何が違うのか

ここまで書いてきて、私は本書のタイトルに対する答えを、かなり明確に言えるようになりました。

AIがなくても、良い上司は良い質問をできます。研修もできます。1on1もできます。サーベイもできます。人事評価もできます。チームミーティングもできます。

一つひとつは、AIがなくてもできます。では、何が違うのでしょうか。

全社員について、毎週、前回の対話を踏まえて問いを返す。過去と現在を比較する。一人ひとりに合わせて質問を変える。本人の行動、本人の受け止め、上司から見た成果を別々に観測する。複数人の共通するパターンを探す。本人だけでは変えられない問題を上司、チーム、組織へ上げる。そして、変えた結果をまた観測する。

これを、人間だけで継続することは極めて難しかった。

だから、AIによって究極的に変わったのは、「人材育成に新しい理論が生まれたこと」ではありません。

人間の時間と記憶の限界によって、断片的にしか行えなかった学習を、全社員について、継続的な経営の循環として運転できるようになったことです。

研修が、単発のイベントではなくなる。1on1が、単発の対話ではなくなる。サーベイが、年に一度の点数測定ではなくなる。人事評価が、半年に一度の判定ではなくなる。すべてが次の行動につながり、次の観測につながり、次のPlanへ戻る。

これまで「点」で存在していた人材育成が、「線」になり、そして「循環」になったのです。

「簡単」になったのは、学習を運転し続けること

本書のタイトルは、『学習する組織は、AIで簡単につくれた！』です。

しかし、「学習する組織」そのものが簡単になったわけではありません。

人を理解することは簡単ではありません。本人の未来と会社の未来をつなぐことも、良いチームをつくることも、正しい評価をすることも、会社の前提を変えることも簡単ではありません。

AIが、それらの答えを代わりに出してくれるわけでもありません。では、何が簡単になったのでしょうか。

学習を、止めずに回し続けることです。

人間だけでは「時間」がなくてできなかった問いかけを、毎週続けられる。人間だけでは「覚えきれなかった過去」を、現在と比較できる。一つの点数では「見えなかった人の状態」を、複数の窓から見ることができる。「一人の問題だと思っていた」ものを、チームや組織の構造として見るができる。そして、それらをPlan、Do、Check、Actとして「つなぎ続ける」ことができる。

AIによって、「学習する組織」を運転するコストを大きく下げられるようになった。これが、私が「AIで簡単につくれた」と言う、本当の意味です。

ただし、だからこそ Human in the Loop が必要になります。

AI が問い、観測し、記憶し、比較し、構造化する。しかし、何になりたいのかを選ぶのは本人です。会社が何を指すのかを決めるのは経営者です。どの Role を任せるのか、何を成果と見るのか、どの Gap に意味があるのか、本人を変えるのか会社を変えるのか、成功や失敗から何を会社の知識へ戻すのかを判断するのは人間です。

AI が問い、観測し、比較し、整理する。人間が意味づけし、選び、判断し、責任を引き受ける。

AI が入ったから、人間が不要になったわけではありません。むしろ AI ができることが増えたからこそ、人間が引き受けるべき意味と責任が、以前よりはっきりしたのだと思います。

完成させない会社

学習する組織(Learning Organization)には、完成形がありません。

会社の仕組みがどれほど良くなっても、環境が変われば、昨日の正解が今日も正しいとは限りません。人も変わります。顧客も変わります。技術も変わります。AI も変わります。

つまり、会社の仕組みだけが完成したまま止まることはできません。

EvaliA も完成品ではありません。問い方が適切でないと分かれば、問いを変える。判定基準が特定の職種に合わない分かれば、ものさしそのものを変える。CAGE 適応モデルについても、実証を重ねる中で見直すべき点が見つかるかもしれません。

これが PDCA です。

完成した瞬間に学習を止めるのであれば、それはもう学習する組織ではありません。人も未完成です。会社も未完成です。AI も未完成です。

継続的に観測する。問い直す。試す。そして、また見る。必要なら、自分たちが正しいと思っていた前提そのものから変える。

学習する組織とは、正解を持っている会社ではなく、正解が変わったときに、自分たちも変わることのできる会社なのだと思います。

二十年前の問いに、いま答える

二十年前の私は、「どうすれば社員を育てられるのか」と考えていました。

その問いが間違っていたとは思いません。ただ、今の私には、その問いだけでは足りません。

社員を育てるだけでは、会社は学びません。仕組みをつくるだけでは、人は学びません。AIを導入するだけでも、学習する組織にはなりません。

本人が未来を考える。会社との接点を見つける。仕事の中で小さく試す。AIとの対話によって振り返る。上司も自分の見方を問い直す。チームが自分たちの働き方を学ぶ。構造に問題があれば、会社が変わる。成功も失敗も、次の人と次の仕組みへ渡す。

そして、変わった仕組みから、また人が学ぶ。

人を育てるのでも、仕組みだけを育てるのでもない。人が仕組みから学び、仕組みも人から学ぶ。

二十年前に描いた二つのベクトルは、今も私の中にあります。ただし、今は、その二つを一度そろえて終わりだとは思いません。

本人が向きたい方向も変わります。会社が向きたい方向も変わります。Roleも、チームも、市場も変わります。だから必要なのは、一度の「一致」ではありません。

循環です。

ずれたら観測する。問い直す。調整する。試す。そして、また見る。AIによって、この循環を止めずに回すことが、初めて現実的になりました。今の私の問いは、二十年前とは少し違います。

「どうすれば、人も会社も、学ぶことを止めずにいられるのか？」

この問いに、完成した答えはありません。

AIも完成した答えを持っていません。私も持っていません。だからこそ、問い続けるのだと思います。

人を変えることから、人と会社がともに学ぶことへ。

そして、私自身もまた、その Human in the Loop の中で、学び続けていきたいと思います。

**東京コンサルティンググループ
会長 公認会計士 久野康成**

[参考文献]

- ・『学習する組織』ピーター・M・センゲ 英治出版、2011 年
- ・『知識創造企業（新装版）』野中 郁次郎（著）、竹内 弘高（著）、梅本 勝博（翻訳）東洋経済新報社、2020 年
- ・『できる若者は3年で辞める！伸びる会社はできる人よりネクストリーダーを育てる』久野康成 出版文化社、2007 年
- ・『やっぱり仕組み化』久野康成 TCG 出版、2024 年
- ・『社長なしで、AI だけで経営やってみた！経営を AI にどこまで任せることが出来たのか？』久野康成 TCG 出版、2026 年
- ・『企業はなぜ進化を止めるのか？The PentaLoop Method 認知を進化させ自ら成長し続ける企業のつくり方』久野康成 TCG 出版、2026 年
- ・『人はなぜ成長を止めるのか？The CAGE Adaptation Model 成功のトラップを超え、自分の人生を選び直す』久野康成 TCG 出版、2026 年
- ・『なぜ人事評価だけでは人は育たないのか』久野康成 TCG 発行小冊子、2026 年

【著者プロフィール】

久野康成 くの・やすなり

久野康成公認会計士事務所 所長

東京コンサルティンググループ創業者

株式会社東京コンサルティングファーム 取締役会長

東京税理士法人 統括代表社員

公認会計士 税理士

26 か国で日系企業の経営支援を行い、組織改革・海外進出・事業承継・AI 経営をテーマに活動。SECI モデル、トヨタウェイ、学習する組織を AI 時代へ実装する新しいマネジメントモデルを提唱。

1965 年生まれ。愛知県出身。1989 年滋賀大学経済学部卒業。

1990 年青山監査法人 プライスウオーターハウス（現・PwC Japan 有限責任監査法人）入所。監査部門、中堅企業経営支援部門にて、主に株式公開コンサルティング業務にかかわる。

クライアントの真のニーズは「成長をサポートすること」にあるという思いから監査法人での事業の限界を感じ、1998 年久野康成公認会計士事務所を設立。営業コンサルティング、IPO コンサルティングを主に行う。

現在、東京、名古屋、大阪、インド、中国、香港、タイ、インドネシア、ベトナム、メキシコほか世界 26 か国にて経営コンサルティング、人事評価制度設計および運用サポート、海外子会社設立支援、内部監査支援、連結決算早期化支援、M&A コンサルティング、研修コンサルティングなど幅広い業務を展開。グループ総社員数約 360 名。

著書に『できる若者は3年で辞める！伸びる会社はできる人よりネクストリーダーを育てる』『もし、かけだしカウンセラーが経営コンサルタントになったら』（以上、出版文化社）『あなたの会社を永続させる方法』（あさ出版）『国際ビジネス・海外赴任で成功するための賢者からの三つの教え 今始まる、あなたのヒーローズ・ジャーニー』『海外直接投資の実務シリーズ』『「大きな会社」ではなく「強い会社」を作る 中小企業のための戦略策定ノート』『創業 26 年 たった 1 人で始めた会計事務所が純キャッシュ 31 億円を貯めた仕組み お金の戦略』『やっぱり仕組み化』『ペントループ』『企業はなぜ進化を止めるのか？』『人はなぜ成長を止めるのか？』（以上、TCG 出版）『現金 10 億経営会社をキャッシュリッチに変える経営戦略』『経営参謀型税理士の使い方 税理士 3.0』（幻冬舎）ほか多数。

