

Le Break-even différé du client : une asymétrie structurelle des modèles FDE

Pourquoi le fournisseur et le client n'observent pas leur retour dans la même fenêtre

Introduction

À l'été 2026, plusieurs grands fournisseurs d'infrastructure et de modèles ont annoncé ou renforcé des dispositifs d'ingénieurs déployés au contact des clients. Les montants engagés se comptent en centaines de millions à quelques milliards de dollars, les effectifs annoncés en milliers d'ingénieurs. Le *forward deployed engineer*, dont un éditeur de logiciels d'analyse a fait il y a une vingtaine d'années un modèle distinctif au service d'agences publiques, est devenu l'un des modèles industriels les plus visibles du déploiement de l'IA en entreprise.

Le discours ambiant traite ce dispositif comme un modèle de *delivery*, une manière plus intime de mettre un système en production. C'est une lecture de surface. Le vrai sujet n'est pas la typologie des FDE. C'est une asymétrie : le fournisseur et le client n'observent pas leur retour dans les mêmes fenêtres temporelles. La question qui décide de la valeur n'est donc pas *comment* on déploie, mais qui finance l'ingénieur déployé, par quel mécanisme son coût est récupéré, et à quel moment le retour net du client peut être établi.

Un point de vocabulaire doit être réglé d'emblée. Par commodité, on parle de ROI. En réalité, un fournisseur reconnaît d'abord un revenu, une marge, un apprentissage ou une option stratégique. Le ROI n'apparaît qu'une fois cette valeur rapportée au coût complet de l'intervention. La confusion entre *retour* et *ROI* n'est pas neutre. Elle permet de présenter comme rendement ce qui n'est encore qu'encaissement.

Nous proposons ici un cadre général, le **Dual Value Horizon Model (DVHM)**, dont le FDE n'est que le premier terrain d'application. Le DVHM décrit une propriété des systèmes sociotechniques dans lesquels la capture de valeur du fournisseur et la démonstration de valeur du client obéissent à des horizons temporels distincts et désynchronisés. Les mécanismes de financement du FDE ne sont pas l'objet du modèle. Ils en sont l'illustration la plus lisible.

La thèse tient en deux phrases. Le mécanisme de récupération fixe la fenêtre dans laquelle le fournisseur peut reconnaître son retour. Le retour net du client relève d'une autre fenêtre, indépendante et le plus souvent plus tardive, car il ne peut être établi

qu'après observation du *coût de permanence*, c'est-à-dire du coût d'exploitation, de gouvernance, de supervision, d'adaptation et de réversibilité sur tout l'horizon, et après attribution causale des bénéfices au déploiement.

La raison en est structurelle, et tient à l'incomplétude des contrats développée plus bas. Le client conserve les droits de contrôle dont dépend la persistance du résultat, de sorte que son retour ne se révèle qu'à l'usage. Le domaine de validité est explicite. Le retour désigne ici le retour financier net après coût complet. Le périmètre est celui du FDE de plateforme, une équipe technique et produit intervenant au contact opérationnel du client pour un fournisseur dont le retour dépend de l'adoption de son produit ou de son écosystème. Cela exclut le conseil indépendant sans intérêt produit.

Le FDE n'est pas seulement un modèle de delivery

Trois questions précèdent toute discussion sur la valeur. Qui règle ou emploie l'équipe déployée. Sur quel acteur le coût retombe *in fine*. Qui est exposé à la perte si la mission échoue. On peut les décliner en six variables de lecture, reprises plus bas en synthèse : payeur immédiat, porteur économique, mécanisme de récupération, actif capturé, horizon du retour et exposition à l'échec.

Quatre mécanismes, et deux métiers

Quatre mécanismes de récupération coexistent. Ils n'engagent ni le même payeur ni le même actif capturé. Ils diffèrent de registre, de financement, de gouvernance, de structure juridique. Mais ils se lisent tous à travers une même variable, la manière dont le fournisseur récupère le coût de l'intervention, car c'est elle qui fixe la fenêtre où il peut reconnaître son retour.

Le premier finance le FDE comme un coût produit, proche d'une quasi-R&D. Le fournisseur porte comptablement l'intervention parce que les apprentissages remontent dans la plateforme. Le terrain client devient un environnement d'apprentissage produit partagé, dont la valeur est principalement capitalisée par l'éditeur, même si le client reçoit adaptation, intégration et résolution accélérée. La contrepartie est une plateforme intégrée qui réduit le coût de coordination tout en augmentant le coût de sortie. On ne peut dire, faute de données publiques comparables, lequel des deux l'emporte. Un éditeur d'analyse dont la part de services a reculé sur plusieurs exercices en fournit une illustration compatible, sans à lui seul établir le mécanisme.

Le deuxième finance le FDE comme un accélérateur de consommation ou de licence. Un engagement de consommation cloud décroît sur l'usage effectif et se perd s'il n'est pas consommé en fin de terme. L'ingénieur déployé peut alors contribuer à convertir un engagement déjà contracté en usages effectifs. Sa fonction opérationnelle ne se réduit pas à ce mécanisme financier. Il peut aussi accélérer le *time-to-value* et

réduire les échecs d'adoption. Le mécanisme économique du fournisseur et la fonction rendue au client sont deux lectures distinctes du même dispositif.

Le troisième finance le FDE par une structure dédiée. Un véhicule d'investissement porte le déploiement et distribue un rendement à ses investisseurs. C'est le cas le plus financiarisé. Il oblige surtout à ne pas confondre quatre acteurs : le rendement de l'investisseur, le retour du véhicule, le bénéfice du fournisseur technologique et la valeur captée par le client ne coïncident pas.

Le quatrième finance le FDE comme une activité de services facturée. Le client paie directement l'équipe ou le package d'implémentation. Ce n'est pas du conseil rebaptisé dès lors que le retour du fournisseur reste attaché à l'adoption de son produit. Le critère n'est pas la facturation. C'est l'intérêt produit, c'est-à-dire ici la stimulation de la consommation du produit sous-jacent.

Ces familles ne classent pas les fournisseurs. Elles décomposent des mécanismes qu'une même mission peut cumuler. Lues selon les six variables, elles se résument ainsi.

Mécanisme	Payeur immédiat	Porteur économique	Actif capturé	Horizon	Exposition principale
Quasi-R&D	Fournisseur	Fournisseur puis clients futurs	Apprentissage produit	Moyen terme	Fournisseur
Consommation ou licence	Fournisseur ou contrat global	Client via la consommation	Usage et engagement	Court à moyen terme	Partagée
Structure dédiée	Investisseurs et véhicule	Véhicule et clients	Distribution, participation, rendement	Moyen à long terme	Investisseurs et clients
Services facturés	Client	Client	Revenu de service et adoption	Court terme	Client, selon contrat

Un cinquième acteur traverse ces familles sans en constituer une : l'intégrateur. Il co-finance parfois une partie du déploiement, capture sa propre marge, et surtout déplace le porteur du coût de permanence. Selon le montage, l'intégrateur porte durablement la capacité par un contrat de services managés, ce qui la rend disponible sans être interne. Ou il la retire avec la fin de son mandat, laissant le client sans la compétence qu'il croyait acquise.

Le fournisseur, d'ailleurs, est rarement un acteur unique. Un éditeur, sa plateforme cloud, une forge logicielle, un fournisseur de modèles, un intégrateur et une *practice* de services peuvent se partager la création de valeur. Il faut alors raisonner en *chaîne de capture* plutôt qu'en fournisseur unique. Chaque maillon a sa propre fenêtre de reconnaissance, ce qui ne réduit pas l'asymétrie mais la démultiplie.

Une seconde distinction, orthogonale au financement, décide de la part du bénéfice réellement attribuable au déploiement, car tous les FDE ne font pas le même métier. Un premier profil analyse les données du client pour en reconstruire l'ontologie et y découvrir des risques ; son terrain est la donnée. Un second profil suppose une connaissance de l'industrie du client, de son paysage concurrentiel, de sa maturité numérique, de ses processus et de ses flux, en sus des compétences techniques de déploiement, pour proposer une transformation d'ampleur. Ce profil est rare.

Un troisième cas mérite d'être isolé sans être surinvesti. Un profil qui maîtriserait de surcroît les aspects réglementaires et juridiques du déploiement rendrait au client une capacité que celui-ci ne détient pas toujours en interne. Ce cas est structurellement peu probable dans le mécanisme de consommation : les incitations économiques du fournisseur ne sont pas alignées avec une maximisation des contraintes juridiques susceptibles de ralentir l'adoption. Ce n'est pas un procès d'intention. C'est une lecture d'alignement d'incitations : le mécanisme de récupération sélectionne le profil, et un profil optimisé pour l'adoption n'est pas optimisé pour la friction réglementaire.

L'observation de terrain confirme cette lecture sans la prouver. Dans plusieurs organisations de services d'éditeurs, la compétence dominante est technique ou d'avant-vente, non une connaissance métier profonde ; les conseillers sectoriels relèvent souvent du *consultative selling*, non de l'architecture d'entreprise appliquée à l'IA. Cet écart n'est pas un détail de gestion des ressources. Lorsqu'un déploiement censé produire une transformation sectorielle est confié à un profil d'avant-vente, le bénéfice réellement attribuable au FDE diminue, et la capacité transférée au client avec lui.

Pourquoi ces distinctions importent-elles. Parce que chaque mécanisme fixe une fenêtre différente dans laquelle le fournisseur peut prétendre avoir créé de la valeur. Et parce que chaque profil décide de la part de cette valeur qui reviendra réellement au client.

Quatre questions sur le retour

Le mot *retour* recouvre quatre propriétés distinctes, qu'il vaut mieux nommer pour les réutiliser : reconnaissance, réalisation, persistance, attribution. Le retour est-il *reconnu* dans les métriques du fournisseur, revenu, marge, valeur actuelle nette estimée. Est-il *réalisé*, encaissé ou coût évité. *Persiste-t-il*, après le retrait de l'équipe et l'intégration des coûts de gouvernance. Est-il enfin *attribuable*, réellement causé par le FDE, établi contre un scénario de référence.

L'attribution n'est pas une quatrième étape après la persistance. C'est une propriété causale qui traverse les trois autres et peut être interrogée à tout moment. Un gain peut être reconnu, réalisé et durable sans être causé par le FDE, simplement parce qu'il se serait produit de toute façon. Un retour encaissé n'est pas encore un retour attribuable. Et sans scénario de référence, il n'y a pas d'attribution du tout : toute démonstration de

ROI sans contrefactuel est un récit, pas une mesure. Cette exigence protège l'analyse contre les *business cases* avant-après qui attribuent au déploiement toute évolution observée.

Ces quatre lectures valent pour chaque acteur. C'est ce qui ruine l'opposition trop commode entre un fournisseur qui gagnerait tôt et un client qui gagnerait tard. La thèse ne dit pas que le fournisseur reconnaît toujours son retour en premier. Un éditeur d'analyse peut le reconnaître très tard ; un véhicule dédié en fonction d'une consommation future ; un hyperscaler plusieurs années après son engagement. Elle dit que les fenêtres de reconnaissance sont *indépendantes*, chacune commandée par son propre mécanisme de récupération. La configuration la plus fréquente reste celle où le fournisseur reconnaît une partie de son retour avant que le client ait pu observer son coût complet.

Deux fonctions de valeur, deux horizons temporels

L'asymétrie entre fournisseur et client ne tient pas seulement à leurs intérêts économiques. Elle tient au fait qu'ils n'observent pas leur retour sur le même horizon. Le fournisseur cherche à capturer rapidement la valeur créée par le déploiement. Le client ne peut établir son retour net qu'après avoir supporté le coût complet de l'exploitation.

L'analyse précédente se généralise sous la forme d'un modèle simple décrivant la désynchronisation entre la capture de valeur du fournisseur et la démonstration de valeur du client. C'est le DVHM annoncé en introduction. Il ne décrit pas un cas particulier de FDE, mais une propriété des systèmes sociotechniques dans lesquels création, reconnaissance et démonstration de la valeur obéissent à des horizons distincts.

La valeur capturée par le fournisseur est représentée par une fonction $CVF(t)$, qui agrège les formes de valeur qu'il reconnaît au cours du temps.

$$CVF(t) = REV(t) + LRN(t) + ECO(t)$$

où $REV(t)$ désigne les revenus effectivement reconnus (licences, abonnements, prestations, consommation contractualisée), $LRN(t)$ la valeur économique des apprentissages réutilisables au bénéfice du fournisseur (amélioration de la plateforme, accélération des futurs déploiements, réduction des coûts de développement, avantages concurrentiels), et $ECO(t)$ la valeur attendue des usages futurs induits par le déploiement (consommation cloud, jetons, licences additionnelles, ventes croisées, réduction du *churn*). Cette fonction n'est pas une identité comptable. C'est une décomposition économique des principaux mécanismes de capture, exprimée dans une unité commune indépendamment du mode de reconnaissance financière.

La dynamique de cette fonction éclaire toute la thèse. $CVF(t)$ croît fortement autour du déploiement, puis sa pente s'effondre. La raison est économique, non mathématique. Chacune de ses composantes voit sa contribution marginale imputable à la présence de l'ingénieur déployé décroître une fois le système en production. Le revenu contractuel est reconnu ou sécurisé. L'engagement de consommation est déjà contracté. Les apprentissages exploitables de la mission sont déjà remontés dans la plateforme. La valeur d'écosystème continue de courir, mais elle ne dépend plus de la présence sur site. Autrement dit, ce n'est pas la valeur qui cesse ; c'est la part de valeur attribuable au FDE qui sature.

Il existe donc un horizon T_F , généralement proche de la mise en production, au-delà duquel la valeur marginale directement imputable au FDE devient faible.

$$T_F = \inf \left\{ t \mid \frac{dCVF(t)}{dt} \approx 0 \right\}$$

T_F marque le point de saturation : l'instant où la dérivée s'annule non parce que le fournisseur cesse de gagner, mais parce que l'équipe déployée n'apporte plus de valeur marginale. L'essentiel de ce que le fournisseur retire du déploiement est déjà reconnu ou sécurisé lorsque l'équipe se retire. T_F n'est pas une constante universelle. Il dépend du mécanisme de récupération retenu. La formalisation ne cherche pas à fixer sa valeur, mais à distinguer sa logique de celle du break-even client.

La trajectoire du client est différente. Sa valeur nette commence négativement, puisqu'il supporte immédiatement le coût du projet et de la constitution de capacité, puis évolue sous l'effet des bénéfices réellement attribuables et du coût de permanence.

$$NVC(t) = -(P + IC_0) + \int_0^t [BA(\tau, C(\tau)) - O(\tau) - K(\tau, C(\tau)) - X(\tau)] d\tau$$

où $BA(\tau, C)$ désigne les bénéfices économiques causalement attribuables au système, établis contre un scénario contrefactuel crédible ; $O(\tau)$ les coûts récurrents d'exploitation technique ; $K(\tau, C)$ les coûts de gouvernance, de supervision humaine, de conformité et d'auditabilité ; $X(\tau)$ les coûts d'adaptation, de remédiation, de réversibilité et de sortie ; P le coût initial du projet ; IC_0 le coût de constitution de la capacité organisationnelle permettant d'exploiter durablement le système.

La nomenclature est tenue stable d'un bout à l'autre. En particulier, la réversibilité est notée $X(\tau)$ et non $R(\tau)$, afin d'éviter toute collision avec le revenu fournisseur $REV(t)$. Un symbole, un sens.

La lecture du signe de l'intégrande explique le retard du break-even. Au début, BA est faible : l'adoption monte, les bénéfices ne sont ni pleinement réalisés ni encore

attribuables contre un contrefactuel. Au même moment, $O + K + X$ est élevé : la gouvernance, la supervision et l'adaptation pèsent surtout dans les premiers temps. L'intégrande est donc négatif ou faiblement positif, et $NVC(t)$ reste sous zéro. Plus tard, si les bénéfices se matérialisent et que les coûts récurrents se stabilisent, l'intégrande passe positif et la valeur nette cumulée finit par croiser l'axe. Ce croisement est tardif par construction, car il ne peut se produire qu'après observation du coût complet de permanence.

Le terme $C(t)$ est la variable médiatrice entre les deux horizons. Il désigne la capacité effectivement transférée et institutionnalisée chez le client, $C(t) \in [0, 1]$. Il relie le DVHM aux travaux antérieurs sur la transférabilité de capacité. Son effet est double et de signe opposé sur les deux termes qu'il gouverne.

$$\frac{\partial BA}{\partial C} > 0 \text{ et } \frac{\partial K}{\partial C} < 0$$

Plus la capacité est transférée, plus les bénéfices attribuables augmentent, car le client sait exploiter, corriger et étendre le système. Et plus les coûts de gouvernance diminuent, car la supervision cesse de dépendre d'un opérateur externe. Un déploiement à faible C maintient BA bas et K haut, retarde le croisement, et peut le rendre inatteignable. Un déploiement à fort C rapproche le break-even. La capacité transférée n'est donc pas une variable de confort. C'est le canal par lequel un déploiement à horizon fournisseur court peut, ou non, produire un retour client dans un horizon fini.

La trajectoire de valeur nette définit alors un temps caractéristique, le break-even du client, premier instant où la valeur nette cumulée devient positive.

$$BET_C = \inf\{t \mid NVC(t) > 0\}$$

La propriété centrale du modèle est une inégalité entre deux durées.

$$BET_C \gg T_F$$

$NVC(t)$ ne calcule pas un ROI mais un temps caractéristique. En analyse dimensionnelle, la différence est fondamentale. Le ROI est une grandeur adimensionnelle ; le break-even est une durée. La question n'est donc plus « quel est le ROI du projet », mais « à partir de quel moment peut-on démontrer que la valeur nette cumulée devient durablement positive ».

Le fournisseur atteint son horizon de reconnaissance alors que le client n'a pas encore observé le coût complet de permanence. Les deux acteurs n'évaluent ni le même objet économique, ni le même horizon temporel. Une adoption élevée au moment de la mise en production valide principalement la fonction de capture du fournisseur. Elle ne

démontre pas encore l'existence d'un retour net durable pour le client. Cette formulation ne dépend plus du cas particulier des FDE. Elle décrit la désynchronisation structurelle entre la fenêtre de capture du fournisseur et la fenêtre de démonstration du client. C'est cette désynchronisation, davantage que le modèle de *delivery* lui-même, qui constitue le cœur économique de la thèse.

Après le retrait du FDE

Le coût de permanence n'est pas seulement technique. Il dépend de l'organisation qui conserve, après le retrait du FDE, la capacité de décider, d'opérer, de financer et de répondre. Ces quatre fonctions, qu'il vaut mieux nommer par leurs verbes plutôt que par le mot-valise *ownership*, ne suivent pas le même acteur. Décider, c'est détenir l'autorité. Opérer, exécuter et maintenir. Financer, prendre en charge les coûts initiaux et récurrents comme l'exposition à l'échec. Répondre, justifier et corriger.

La théorie des contrats incomplets, de Grossman et Hart à Hart et Moore, fournit le mécanisme. Lorsque les états futurs d'un système ne peuvent pas tous être spécifiés au contrat, ce sont les droits de contrôle résiduels qui déterminent qui décide dans les situations non prévues. C'est parce que ces droits restent chez le client que la fenêtre de son retour est repoussée à l'usage. Le contrôle détermine qui peut agir ; le contrat organise qui doit payer. Le verbe importe, car le porteur final du coût dépend aussi de la loi, de la responsabilité civile, de l'assurance et de la solvabilité des acteurs.

Une décision prise par un système automatisé crée une fonction d'imputabilité qui doit rester assignée à une personne ou une institution identifiable. Elle mobilise une capacité durable de supervision, de justification et de remédiation. C'est elle qui alimente le coût de gouvernance K de l'équation, et c'est elle que la transférabilité C vient alléger. Il serait tentant d'en conclure qu'il manque un responsable unique du résultat. Ce serait une erreur. La performance d'un système sociotechnique est légitimement collective, et désigner un propriétaire unique reviendrait à faire porter un risque à un acteur sans contrôle suffisant. L'enjeu est de pouvoir reconstituer qui contrôlait quoi, qui devait intervenir et selon quelle règle lorsque le système a dérivé. Autrement dit, l'allocation des responsabilités doit rester reconstructible.

La responsabilité qui reste au client n'est soutenable qu'à une condition, et cette condition n'est pas l'internalisation. Trois notions se succèdent ici sans se confondre. La *compétence* est un savoir individuel. La *capacité* une aptitude organisée à agir. La *capacité institutionnalisée* une capacité dotée de rôles, d'autorité et de budget durables. Le bon critère n'est pas que la capacité soit interne, mais qu'elle soit durablement disponible, interne ou contractuellement garantie par un opérateur managé identifiable. Encore faut-il ne pas confondre les degrés de ce transfert, ceux-là mêmes que $C(t)$ agrège. Les moyens peuvent être livrés sans que la compétence soit transférée. La compétence peut être transférée sans être institutionnalisée. La capacité peut être

institutionnalisée sans être réellement exercée. Une compétence documentée n'est pas encore une capacité institutionnelle.

Contre-modèles, limites et conclusion

La contre-thèse la plus sérieuse est le contrat au résultat. On objecte qu'un *gainshare* bien conçu alignerait le fournisseur sur la performance du client. L'objection est réelle, mais elle bute sur la structure du système. Même bien conçu, le *gainshare* porte sur une performance coproduite par un contrôle distribué. Le fournisseur contrôle principalement certains composants techniques, les conditions d'accès à la plateforme et une partie des remédiations. Le client contrôle principalement les données métier, les décisions opérationnelles, l'adoption et les seuils. Le problème ne tient pas d'abord à la mauvaise foi des acteurs, mais à la non-contractibilité d'une performance dont aucun ne contrôle ni n'observe seul toutes les variables.

Restent les limites propres au modèle, qu'il vaut mieux énoncer que dissimuler.

Le DVHM est **bilatéral**. Il oppose valeur fournisseur et valeur client, et cette opposition est aussi son angle mort. Certains déploiements créent une valeur réglementaire, écologique ou systémique qui n'est captée par aucun des deux acteurs et qui n'apparaît dans aucune des deux fonctions. Un déploiement peut détruire du retour client tout en produisant un bénéfice collectif, ou l'inverse. Le modèle ne le voit pas. Il faudrait, pour l'étendre, une troisième fonction de valeur sociale, dont les termes ne se réduisent ni à *CVF* ni à *NVC*. Cette extension n'est pas traitée ici.

Le coût de permanence est traité comme un **coût**, jamais comme un **actif**. C'est une simplification assumée. La gouvernance, la documentation, la montée en compétence et l'auditabilité produites pendant l'exploitation créent aussi des options futures : réutilisation sur d'autres systèmes, conformité déjà acquise, capacité de négociation accrue. Il existe donc une *valeur de permanence* symétrique du coût de permanence, que le modèle ne capture pas. La formaliser reviendrait à ajouter à l'intégrande un terme positif dépendant lui aussi de *C*. C'est un prolongement identifié, non un acquis.

La fonction $C(t)$ elle-même reste sous-spécifiée. Le modèle pose son existence, ses deux dérivées partielles et son rôle médiateur. Il ne propose ni forme fonctionnelle, ni protocole de mesure. Or c'est probablement la variable la plus déterminante et la plus difficile à observer. La mesurer suppose de distinguer la compétence livrée de la capacité exercée, ce qu'aucun indicateur d'adoption ne fait.

Enfin, une limite de démonstration doit être nommée sans détour. Ce texte établit que les fenêtres de reconnaissance et de démonstration sont *distinctes* et gouvernées par des mécanismes différents. Il n'établit pas encore que cette différence *cause* effectivement l'écart de retour observé entre fournisseur et client. Il montre une corrélation conceptuelle, pas une causalité mesurée. La distinction est exactement celle que le

modèle impose par ailleurs à ses propres objets : un lien reconnu n'est pas un lien attribué.

Cette limite dessine le programme de recherche plutôt qu'elle ne l'invalide. L'hypothèse testable est la suivante : dans une proportion substantielle de missions, l'inclusion du coût de permanence et l'attribution causale modifient matériellement l'évaluation initiale du retour client, et ce retour ne peut être établi de façon fiable dans la même fenêtre que celui du fournisseur. Le protocole tient en quelques lignes. Sur une population d'une cinquantaine de missions FDE, comparer le *business case* initial au break-even observé, puis au retour net observé à vingt-quatre mois. Documenter pour chacune le coût de projet, le coût de permanence par ses cinq catégories, le niveau d'adoption, la capacité effectivement transférée C , et un scénario contrefactuel. Le test discriminant porte sur C : si les missions à fort C atteignent BET_C significativement plus tôt, à mécanisme de récupération comparable, la médiation cesse d'être narrative et devient mesurée. Une hypothèse philosophique devient alors une hypothèse falsifiable, et le DVHM un cadre réfutable plutôt qu'une grille de lecture.

Les contre-modèles réels ne confirment pas la thèse : ils circonscrivent son domaine de validité. Un opérateur managé rémunéré sur une performance persistante, un contrat allouant de façon reconstructible les responsabilités partielles, un transfert de capacité réellement institutionnalisé, montrent les conditions qui évitent le problème. La thèse décrit une structure fréquente du FDE de plateforme, non une loi universelle. Elle ne prétend pas que le retour net du client devient généralement négatif. Un déploiement correctement analysé peut afficher un retour client positif après coût complet, et ce cas ne la réfute pas.

Le débat sur les ingénieurs déployés s'est focalisé sur le *delivery*, sur la manière la plus efficace de mettre un système en production. C'était la mauvaise question. Le FDE ne modifie pas seulement le coût du déploiement. Il décale dans le temps les conditions sous lesquelles chacun peut prétendre avoir créé de la valeur. En amont, le mécanisme de récupération fixe la fenêtre où le fournisseur reconnaît son retour. En aval, le coût de permanence et l'attribution causale fixent, dans une autre fenêtre, celle où le client peut établir le sien. Le *business model* organise la fenêtre du retour fournisseur ; le coût de permanence révèle le retour client. Tant que ces deux fenêtres restent confondues, les tableaux de bord d'adoption continueront de prendre des revenus précoces pour des ROI démontrés.

Un dernier mot sur la portée. Le FDE n'est ici que le premier terrain d'application du DVHM. La chaîne qui va du mécanisme de récupération à la fenêtre de reconnaissance, puis à la réalisation, à la persistance, à l'attribution et enfin au retour démontrable, ne dépend pas de la présence d'ingénieurs déployés. Elle vaut pour les ESN en régie longue, les grands programmes ERP, les plateformes cloud, les dispositifs médicaux logiciels, les jumeaux numériques et les transformations organisationnelles. Partout où un fournisseur

reconnaît son retour dans une fenêtre que le client ne partage pas, la même asymétrie opère. Le FDE a seulement le mérite de la rendre visible.

Les modèles FDE ne déplacent pas seulement la distribution de la valeur. Ils déplacent la distribution temporelle de la création de valeur.