



xpdeep

Self-Explainable Deep Learning

**Boostez la compétitivité de vos produits, opérations et services,
avec le seul framework de deep learning auto-explicable**

DEPLOYER DES MODELES D'IA SANS LES COMPRENDRE TOTALEMENT EST UN CHALLENGE STRATEGIQUE



**ENJEU D'ALIGNEMENT
DES PERFORMANCES SUR
LES OBJECTIFS METIERS**



ENJEU DE CONFORMITE

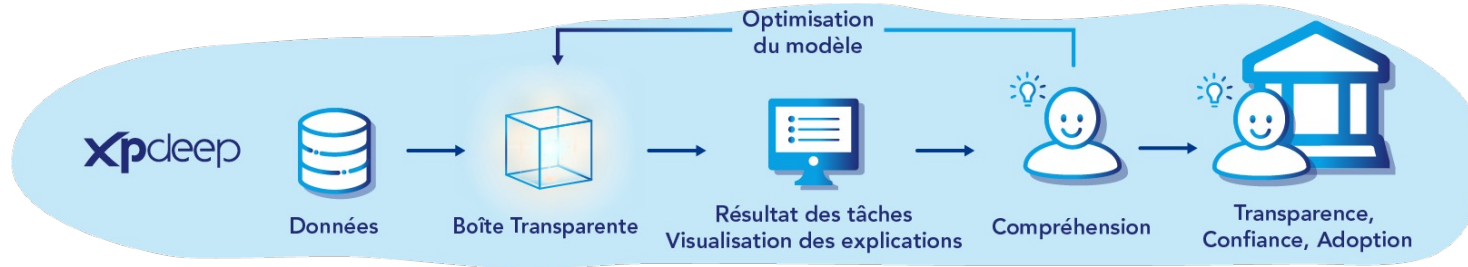


**ENJEU DE
CONFIANCE ET
D'ADOPTION**



**ENJEU DE
GESTION DU RISQUE**

LA SOLUTION XPDEEP : EXPLICABILITÉ, PERFORMANCE ET FRUGALITÉ



CARACTÉRISTIQUES D'XPDEEP

- Conception de modèles auto-explicables sans sacrifier la performance
- Explication de modèles existants, sans altérer leur structure ni leur performance
- Solution pour tout type de tâche basique ou avancée (détection d'anomalies, forecasting, contrefactuel, fairness, classification,...)
- Solution basée sur des approches d'IA les plus adaptées et efficaces sur des données structurées (données capteurs, séries temporelles, séquences, images...)
- Explications accessibles via une interface interactive et flexible

EXPLICATIONS DES MODELES ET DE LEUR FONCTIONNEMENT

1- Expliquer le fonctionnement d'un modèle

- Caractériser les décisions déployées par le modèle et évaluer leurs qualités
- Révéler les facteurs impliqués dans les décisions et leurs importances
- Identifier les parties fortes et faibles (non robustes, sources d'erreurs...) du modèle

→ Bénéfices

- **Accélérer la conception de modèles robustes et performants**
 - Améliorant les décisions non robustes et sources d'erreurs
 - Ajoutant (resp. supprimant) des facteurs utiles (resp. inutiles)
- **Favoriser la conception de modèles frugaux**
 - Adaptant et contrôlant la complexité à attribuer
 - Supprimant les facteurs révélés non contributifs

EXPLICATIONS CAUSALES DES PREDICTIONS/INFERENCE

2- Expliquer les causes des prédictions et inférences

- Caractériser les décisions impliquées dans les prédictions et évaluer leurs qualités
- Révéler les facteurs contributifs des prédictions et leurs importances
- Identifier les décisions et facteurs sources des erreurs de prédiction

→ Bénéfices

- Justifier les erreurs de prédiction, pour des raisons de responsabilité ou de conformité réglementaire
 - *Pourquoi le modèle a échoué dans la détection d'une forme de signal ? Quels capteurs ont été sources d'erreurs, pourquoi ? Comment rectifier ces signaux, à l'entrée ? dans le modèle ? ...*
- Générer des scénarios types (contrefactuels) pour des besoins métiers, d'anticipation de cas indésirables, de calibration de process, ...
 - *explications par le contrefactuel : génération de changements minimaux de données d'entrée, pour qu'un évènement soit bien détecté, un emprunt bancaire soit accepté, une panne soit évitée...*
- Mener des analyses « Fairness/Comparative » confrontant le comportement du modèle sur des groupes différents
 - *le modèle déploie-t-il les mêmes décisions, mêmes facteurs et importances pour l'identification de deux formes de signaux radars ? pour prédire les pannes de pièces provenant de deux fournisseurs différents ?...*

DES EXPLICATIONS INTELLIGIBLE ET ACTIONABLE VIA UNE INTERFACE INTERACTIVE ET FLEXIBLE



Diversité et large éventail d'explications

- Décisions déployées, facteurs et degrés d'importance
- Indicateurs de qualité des explications
- Analyse contrefactuelle multifactorielle
- Analyse Fairness / Comparative



Explications sous diverses formes

- Exemples, contre-exemples
- Résumés textuels
- Graphiques intelligibles



Outil collaboratif

- Dialogue avec l'ingénierie et les utilisateurs métiers
- Débloquer ensemble des points de blocage critiques
- Expliquer aux parties prenantes pour la confiance, la conformité et la gestion du risque

Moments
Eureka



XPDEEP : UNE NOUVELLE FAÇON DE CONCEVOIR ET D'UTILISER LES MODÈLES DEEP



COMMENT UTILISER XPDEEP ?

- Utilisation des bibliothèques XPDeep sous Pytorch
- Disponible on-prem (inc. air-gapped) et sur les clouds publics
- Conception des modèles selon des usages et fonctions standards
- Exporter les modèles dans des formats standards
- Déploiement dans divers environnements

COMMENT OPTIMISER LES MODÈLES ?

- Identifier les sources d'erreurs et d'instabilité des modèles, pour des modèles plus performants et plus robustes, plus rapidement
- Identifier les facteurs d'entrée inutiles, pour des modèles plus frugaux et moins complexes
- Augmenter/diminuer la complexité des zones sous-ajustées/sur-ajustées
- Adaptation flexible des modèles aux contraintes métiers



EXPLIQUER QUOI ET POURQUOI ?

- Expliquer le modèle et ses inférences à des fins d'adoption et de confiance
- Expliquer les inférences générées pour analyser et contrôler les prédictions futures
- Expliquer les prédictions passées pour l'audit et l'examen des scénarios critiques, des incidents, etc.

QUAND EXPLIQUER DEVIENT STRATÉGIQUE ?

- Pour un déploiement en confiance dans les logiciels et machines
- Conformité réglementaire
- Identifier les facteurs de risques à des fins :
 - juridiques,
 - de certification,
 - d'assurance,
 - financières...



A PROPOS DE XPDEEP



Technologie unique au monde / 6 ans de R&D



Premier framework de deep learning européen



Créée en avril 2023, basée à Grenoble



Membre du collectif Confiance AI



Transfert de technologie de l'Université Grenoble Alpes



Membre du collectif Deep Green



17 personnes



Accélérée par HPE



bpifrance

i-Lab



Université Grenoble Alpes

linksium
technology transfer & startup building
Grenoble Alpes



FONDATEURS



AHLAME DOUZAL

CSO

- PhD en ML/data analysis
- Professeure Associée – UGA
- 60 publications sur AI/ML/DL
- Experte de l'IA appliquée aux données temporelles (>20 ans)



STANISLAS CHESNAIS

CEO

- 30 années création et direction de startup dans le BtoB
- Canada, France, USA
- Consulting puis startups
- CEO Netsize 0 to 100M€ de CA net



Self-Explainable Deep Learning

**N'HÉSITEZ PAS À NOUS CONTACTER POUR PLUS
D'INFORMATION**

Nicolas Schu
Responsable Commercial
nschu@xpdeep.com