

ÉPI-ATTENTION : Une attention contextuelle adaptative pour la pertinence dynamique des caractéristiques dans les réseaux de neurones

Faten Chaeib-Chakchouk*, Mohamed Djallel Dilmi*

* Efrei-Paris | Université Panthéon Assas
30-32 Avenue de la République, 94800 Villejuif, France
faten.chakchouk@efrei.fr, djallel.dilmi@efrei.fr

Résumé. Dans cet article, nous introduisons l'Épi-Attention, un mécanisme d'attention contextuelle qui ajuste dynamiquement la pertinence des caractéristiques dans les réseaux de neurones en fonction d'informations externes. Contrairement aux méthodes classiques, l'Épi-Attention permet au modèle de moduler l'importance des caractéristiques en fonction du contexte, améliorant ainsi la capture des propriétés spécifiques aux classes. Nous présentons deux variantes : l'Épi-Attention à Produit Scalaire Réduit et l'Auto-Épi-Attention, qui ajustent la pertinence des caractéristiques selon des informations externes ou internes. L'Épi-Attention améliore l'interprétabilité des modèles, avec des résultats probants sur des jeux de données comme Wisconsin Breast Cancer, Bank Marketing et ABIDE-II.

1 Introduction

Dans des domaines comme la santé et la finance, l'explicabilité des modèles est cruciale pour assurer une prise de décision alignée avec les connaissances des experts. Les modèles traditionnels tels que la régression logistique, les MLPs et les arbres de décision souffrent de certaines limitations : ils appliquent des poids statiques aux caractéristiques, manquant souvent les corrélations locales spécifiques aux classes ou aux sous-populations sous-représentées.

D'autres part, les mécanismes d'attention introduits par Vaswani et al. (2017) sont devenus des piliers de l'apprentissage profond, permettant aux modèles de se concentrer sur les aspects les plus pertinents des données d'entrée Hassanin et al. (2024); Xiao et al. (2024); Bo et al. (2024). Cependant, les mécanismes d'attention globaux, tels que ceux utilisés dans les Transformers, attribuent des poids dynamiques en fonction des relations entre les éléments d'entrée. Bien que cette approche soit puissante pour capturer des dépendances globales, elle peut parfois ne pas saisir pleinement les corrélations locales ou les structures spécifiques, en particulier dans des scénarios impliquant des données hétérogènes, conflictuelles et/ou dépendant d'un contexte. Ces limitations motivent le développement de mécanismes capables d'intégrer des informations contextuelles externes ou internes pour mieux s'adapter aux particularités des données.

Pour répondre à ces défis, nous introduisons l'Épi-Attention, un mécanisme d'attention dynamique et contextuel qui ajuste la pertinence des caractéristiques en fonction des informations internes et externes. Nous proposons également l'Auto-Épi-Attention, qui affine la pertinence des caractéristiques en fonction de la cohérence interne des données, sans avoir besoin de contexte externe. Ces mécanismes améliorent à la fois l'adaptabilité des modèles et leur interprétabilité, tout en capturant des corrélations locales spécifiques aux classes.

2 État de l'art

Les mécanismes d'attention sont essentiels dans les modèles d'apprentissage profond pour des tâches telles que la compréhension automatique, la reconnaissance d'actions et le traitement du langage naturel Britz et al. (2017); Vaswani et al. (2017). Initialement, ces mécanismes se concentraient sur des parties spécifiques des données d'entrée via des vecteurs fixes. Des avancées récentes, comme le réseau Bi-Directional Attention Flow (BIDAF) Seo et al. (2017), intègrent des processus hiérarchiques pour représenter le contexte à différents niveaux. De même, les réseaux Global Context-Aware Attention LSTM Liu et al. (2017) et Liu et al. (2018) améliorent l'attention dans la reconnaissance d'actions 3D en tenant compte des informations contextuelles globales Liu et al. (2021). Dans le traitement du langage naturel, des modèles comme TA-Seq2Seq Xing et al. (2017) et Tri-Attention Yu et al. (2022) incorporent des connaissances externes et modélisent explicitement les interactions entre contextes, requêtes et clés pour améliorer la performance.

Les mécanismes d'attention ont également été appliqués à la reconnaissance des émotions. Par exemple, CAER-Net Lee et al. (2019) utilise les expressions faciales et les informations contextuelles pour améliorer la reconnaissance émotionnelle. Enfin, le Transformer Enrichi en Connaissances (KET) Zhong et al. (2019) exploite des connaissances externes pour enrichir les conversations textuelles.

En résumé, les mécanismes d'attention sont passés de simples approches unidirectionnelles à des méthodes sophistiquées et sensibles au contexte, améliorant les performances dans divers domaines.

3 Notations et définitions

Cette section présente les notations utilisées tout au long de cet article : l'ensemble $\mathcal{D}$ représente l'ensemble des observations. Chaque élément $x \in \mathcal{D}$ correspond à un enregistrement spécifique d'un individu de la population étudiée. Nous considérons également l'ensemble $\mathcal{M}$ qui représente les métadonnées disponibles, appelées contexte, pour chaque élément $x \in \mathcal{D}$, ainsi que la fonction χ qui associe les éléments de $\mathcal{D}$ aux éléments de $\mathcal{M}$.

$$\begin{aligned} \chi: \mathcal{D} &\rightarrow \mathcal{M} \\ x &\mapsto c \end{aligned} \tag{1}$$

Les observations et les métadonnées sont supposées être dans un format séquentiel brut, c'est-à-dire $x = [x^1, x^2, \dots, x^j, \dots, x^p]$ (resp. $c = [c^1, c^2, \dots, c^l]$), avec p (resp. l) la longueur de la séquence x (resp. c), et $x^j, j = 1, \dots, p$ pouvant être de différents types de données.

Nous supposons que nous disposons d'observations correspondant à N individus, représentés par le sous-ensemble $X = \{x_i \mid i = 1, \dots, N\} \subset \mathcal{D}$ et $\mathcal{C} = \{c \in \mathcal{M} \mid c = \chi(x)\} \subset \mathcal{M}$. Bien entendu, il est supposé que X est représentatif de la population étudiée et servira de jeu d'entraînement pour l'estimation des paramètres de divers modèles.

Le traitement de la séquence brute $x \in X$ peut poser des défis, en particulier lorsqu'il s'agit de types de données mixtes (images, mots,...). Plusieurs méthodes visent, dans un premier temps, à encoder/intégrer les informations contenues dans x_i dans un espace vectoriel pour surmonter ces défis de Kok et al. (2024); Sahoo et Chakraborty (2020); elles le font en :

- définissant une fonction d'encodage seq2one f qui encode la séquence complète $x \in X$ sous la forme d'un vecteur $v \in \mathbb{R}^d$:

$$\begin{aligned} f: X &\rightarrow \mathbb{R}^d \\ x &\mapsto v \end{aligned} \quad (2)$$

- ou définissant une fonction d'encodage seq2seq g qui encode chaque élément $x^j \in x$ (avec $x \in X$) sous la forme d'un vecteur $v^j \in \mathbb{R}^d$, ainsi :

$$\begin{aligned} g: X &\rightarrow \mathbb{R}^{d \times p} \\ x &\mapsto [v^j]_{j=1}^p \end{aligned} \quad (3)$$

4 Formulation du problème

Soit $X = \{x_i \mid i = 1, \dots, N\} \subset \mathcal{D}$ l'ensemble de données étudié et $\mathcal{C} = \{c \in \mathcal{M} \mid c = \chi(x)\} \subset \mathcal{M}$ l'ensemble de contexte associé.

En apprentissage automatique, il est courant de traiter l'observation de données $x_i \in X$, c'est-à-dire la séquence brute $x_i = [x_i^1, x_i^2, \dots, x_i^j, \dots, x_i^p]$, en supposant que tous les éléments x_i^j de x_i sont indépendants et pertinents en tant qu'entrées.

Il existe des situations où des informations supplémentaires (contexte), $c_i = \chi(x_i)$, deviennent disponibles. Ces informations externes soulèvent des questions sur la pertinence de certains éléments de x_i en fonction des evidences disponibles. En d'autres termes, connaître certaines informations externes c_i peut augmenter (ou diminuer) la pertinence de certaines variables pour une observation spécifique x_i . Cela nous amène essentiellement à évaluer la plausibilité de l'hypothèse H_i^j liée au $j^{\text{ème}}$ élément de l'observation x_i sur la base des evidences E_i fournies, avec :

- $H_i^j : x_i^j$ est pertinent pour x_i (hypothèse)
- $E_i : c_i$ est une information supplémentaire concernant x_i (preuve/évidence)

$\forall i = 1, \dots, N$ et $j = 1, \dots, p$

Considérons un exemple où nous cherchons à prédire le résultat d'une campagne de marketing bancaire en utilisant plusieurs variables numériques telles que l'âge, le solde, la durée du dernier contact, ... notées respectivement $[x^1, x^2, x^3, \dots]$.

Les experts¹ affirment que la caractéristique x^3 ("durée du dernier contact") "affecte fortement la cible de sortie...", ce qui suggère sa grande pertinence.

1. <https://archive.ics.uci.edu/dataset/222/bank+marketing>

Cependant, la pertinence de x^3 peut être influencée par une information externe supplémentaire c^1 , qui indique si le client a déjà été contacté. Si $c^1 = 1$ ("oui, il a été contacté"), la pertinence de x^3 reste élevée. À l'inverse, si $c^1 = 0$ ("non, il n'a pas été contacté auparavant"), la pertinence de x^3 diminue considérablement. Dans ce cas, un modèle bien entraîné devrait prendre en compte d'autres caractéristiques lorsque la caractéristique x^3 est indéfinie. Ainsi, nous pouvons évaluer la plausibilité de l'hypothèse $H_i^3 : x_i^3 \text{ est pertinente}$ comme élevée sur la base des évidences $E_i : c_i^1 = 1$, tandis que H_i^3 serait jugé peu plausible sur la base des évidences $E_i : c_i^1 = 0$.

Cet exemple illustre comment la plausibilité d'une hypothèse concernant la pertinence d'une caractéristique x^j peut être ajustée en incorporant des informations externes supplémentaires c^i . Par conséquent, un modèle entraîné devrait prêter plus d'attention aux caractéristiques pertinentes pour chaque observation x_i , de manière à s'aligner sur la plausibilité de l'hypothèse « $H : la caractéristique x_i^j est pertinente pour l'observation x_i » pour $j \in \{1, \dots, p\}$ ».$

5 Épi-Attention

5.1 Définition

Soit $X = \{x_i \mid i = 1, \dots, N\} \subset \mathcal{D}$ l'ensemble de données étudié et $\mathcal{C} = \{c \in \mathcal{M} \mid c = \chi(x)\} \subset \mathcal{M}$ l'ensemble de contexte associé.

En considérant à la fois la séquence d'entrée $x_i = [x_i^1, x_i^2, \dots, x_i^j, \dots, x_i^p] \in X$ et le contexte associé $c_i = \chi(x_i)$, nous proposons/définissons l'« Épi-Attention » comme l'ensemble de fonctions $\mathcal{F}$ qui associent les couples $(x_i, c_i) \in X \times \mathcal{C}$ à un vecteur $a_i = [a_i^j]_{j=1}^p \in \mathbb{R}^p$, tel que $\forall f \in \mathcal{F}$:

$$\begin{aligned} f : X \times \mathcal{C} &\rightarrow \mathbb{R}^p \\ (x_i, c_i) &\mapsto a_i = [a_i^j]_{j=1}^p \end{aligned} \quad (4)$$

avec : a_i^j modélisant la plausibilité de l'hypothèse $H_i^j : x_i^j \text{ est pertinente pour } x_i$ en fonction des preuves c_i , $\forall j = 1, \dots, p$ et $\forall i = 1, \dots, N$.

Sur la base de cette définition, nous notons $a_i = f(x_i, c_i)$ le vecteur d'attention. Il est composé de p éléments $a_i^j \in a_i$ qui évaluent la pertinence de $x_i^j \in x_i \forall j = 1, \dots, p$ en fonction du contexte fourni. Ainsi, il nous permet d'attribuer une attention et une importance appropriées aux caractéristiques pertinentes.

Par conséquent, une grande valeur pour a_i^j augmente la pertinence de x_i^j , tandis qu'une petite valeur pour a_i^j réduit son influence, comme illustré dans la figure 2. Un épi-vecteur $\tilde{x}_i = [\tilde{x}_i^1, \tilde{x}_i^2, \dots, \tilde{x}_i^j, \dots, \tilde{x}_i^p]$ est calculé en effectuant une multiplication élément par élément entre le vecteur d'attention a_i et la séquence d'entrée x_i . Cette opération peut être exprimée par : $\forall i = 1, \dots, N, \forall j = 1, \dots, p, \tilde{x}_i^j = a_i^j \times x_i^j$.

En résumé, l'objectif du modèle Épi-Attention est de déterminer les poids d'attention optimaux pour chaque élément x_i^j dans la séquence d'entrée x_i en fonction de leur pertinence selon le contexte c_i . En renforçant ou en atténuant stratégiquement certaines sections de la séquence x_i , le vecteur de poids d'attention a_i joue un rôle crucial dans la formation de la sortie, permettant un contrôle dynamique et précis de la pertinence des différentes caractéristiques.

Il est important de noter que lorsque le traitement de la séquence brute $x_i \in X$ pose des défis, notamment en raison des types de données mixtes (images, mots,...), le concept d'épi-attention reste le même en utilisant l'encodage de séquence approprié comme mentionné dans l'équation (2) et/ou (3). La section suivante présente notre implémentation proposée pour modéliser efficacement un modèle Épi-Attention.

5.2 Épi-Attention à produit scalaire réduit

Pour calculer les poids d'attention $a_i = [a_i^j]_{j=1}^{d_v}$ pour $i = 1, \dots, N$, nous proposons une nouvelle implémentation appelée "Épi-Attention à produit scalaire réduit".

Soit $x \in X$ la séquence d'entrée et c le contexte associé. En utilisant un réseau neuronal net^Q qui associe la séquence d'entrée x à une matrice $Q \in \mathbb{R}^{d_k \times p}$ composée de p vecteurs de requête de dimension d_k :

$$\begin{aligned} net^Q : X &\rightarrow \mathbb{R}^{d_k \times p} \\ x &\mapsto Q \end{aligned} \quad (5)$$

avec $Q = [q_1, \dots, q_j, \dots, q_p]^T$ et $q_j \in \mathbb{R}^{d_k}$ représentant l'encodage de la requête de l'élément $x^j \in x, \forall j = 1, \dots, p$. De même, la séquence de contexte c est mappée dans une matrice $K \in \mathbb{R}^{d_k \times p} = [k_1, \dots, k_j, \dots, k_p]^T$ composée de p clés via un réseau neuronal net^K :

$$\begin{aligned} net^K : X &\rightarrow \mathbb{R}^{d_k \times p} \\ x &\mapsto K \end{aligned} \quad (6)$$

avec $K = [k_1, \dots, k_j, \dots, k_p]^T$ et $k_j \in \mathbb{R}^{d_k}$ représentant l'encodage de la clé pour le contexte de l'élément $x^j \in x, \forall j = 1, \dots, p$.

Pour obtenir l'élément a^j des poids épi-attentionnels $a = (a^j)_{1 \leq j \leq p}$, nous calculons le produit scalaire entre la requête q_j et la clé correspondante k_j , le mettons à l'échelle en le divisant par d_k et appliquons une fonction d'activation (ex. sigmoïde, ...etc.) pour normaliser :

$$a^j = f \left(\frac{q_j \cdot k_j}{d_k} \right) \quad (7)$$

Enfin, lorsque le traitement de la séquence brute x est complexe, un réseau neuronal net^V est utilisé pour intégrer la séquence brute x conformément aux formules 2 ou 3.

Il est important de noter que notre implémentation de « l'Épi-Attention », représentée dans la figure 2, diffère de l'Attention à Produit Scalaire réduit traditionnelle introduite dans Vaswani et al. (2017). Contrairement au mécanisme d'attention conventionnel, qui met principalement l'accent sur la correspondance entre la séquence d'entrée x et la séquence de sortie y en utilisant des éléments de x comme déclencheurs d'attention, l'approche Épi-Attention déplace son attention vers la séquence d'entrée x elle-même. Elle vise à réévaluer la pertinence de ses éléments en incorporant des informations externes c . Ce changement de perspective permet une compréhension et une utilisation plus complètes de la séquence d'entrée au sein du mécanisme d'attention. Par conséquent, au lieu de générer une matrice d'attention, nous produisons

un vecteur de dimension d_v qui correspond au vecteur de représentation v . En outre, cette distinction nous permet de transformer/re-pondérer les données d'entrée dans le même espace de représentation $\mathcal{R}^{d_v}$ et peut conduire à des modèles plus explicables comme montré dans la section 6.

5.3 Auto-Épi-Attention

Le concept d'Épi-Attention est particulièrement utile pour évaluer la cohérence de la séquence d'entrée x . L'Auto Épi-Attention régule l'information transmise par la séquence d'entrée x elle-même, en renforçant les caractéristiques les plus pertinentes tout en minimisant les incohérences. Contrairement à l'auto-attention classique, qui évalue chaque élément isolément, l'Auto Épi-Attention gère les conflits en considérant l'impact de tous les éléments $x - [x^j]$ (tous les éléments sauf x^j) sur la pertinence de x^j dans la séquence x .

Pour approfondir cela, nous examinons l'hypothèse H_i^j liée au j^{eme} élément d'une séquence x_i :

- $H_i^j : x_i^j$ est pertinent pour x_i (hypothèse)
- $E_i : c_i$ est une information interne = x_i (preuve/évidence)

Ainsi, l'Auto Épi-Attention vise à préserver la cohérence interne d'une observation x_i .

6 Expérimentations et Résultats

Cette section présente trois applications où l'Épi-Attention et l'Auto-Épi-Attention peuvent être exploitées pour améliorer l'explicabilité des modèles, ajuster dynamiquement la pertinence des caractéristiques et/ou faciliter le traitement des types de données mixtes. Ces applications répondent à des défis spécifiques à chaque domaine tels que :

App. - Domaine/Dataset	Description
1 - Breast Cancer (Wisconsin)	Auto-Épi-Attention pour ajuster la pertinence des caractéristiques internes, alignée sur les connaissances médicales.
2 - Bank Marketing	Épi-Attention pour pondérer dynamiquement les caractéristiques numériques selon les données qualitatives afin d'améliorer les prédictions client.
3 - Neuroimagerie (Autisme, données multi-sites)	Épi-Attention pour ajuster la pertinence des caractéristiques selon les variations entre sites, pour un modèle plus précis.

TAB. 1 – Applications de l'Épi-Attention dans différents domaines

L'Épi-Attention a été conçue comme une couche de réseau neuronal, facilitant son intégration dans divers modèles existants. Cette approche permet aux chercheurs de bénéficier de ses avantages dans leurs cadres préférés et d'exploiter l'Épi-vecteur amélioré $\tilde{x}$.

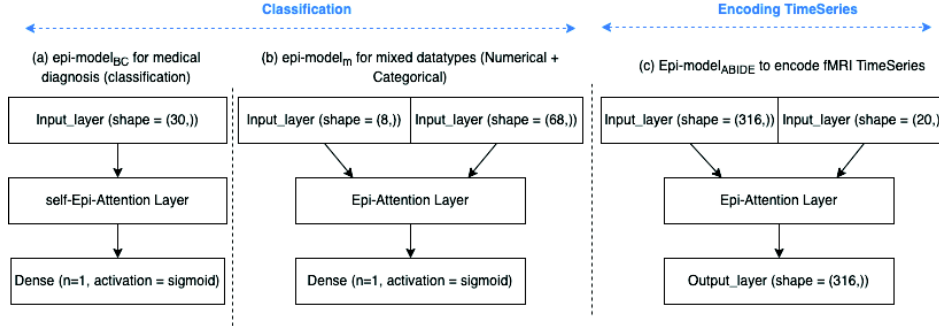


FIG. 1 – Architectures des trois épi-modèles utilisés, de gauche à droite : $epi - model_{BC}$ pour le diagnostic médical avec le Wisconsin Breast Cancer Dataset, $epi - model_m$ pour la gestion des types de données mixtes avec le Bank Marketing Dataset, $epi - model_{ABIDE}$ pour l'analyse de séries temporelles multi-sites avec ABIDE-II.

6.1 Application 1 : Diagnostic Médical (Wisconsin Breast Cancer Dataset)

Nous proposons $epi - model_{BC}$ représenté dans la figure 1 pour le diagnostic médical avec le Wisconsin Breast Cancer dataset, un réseau neuronal conçu pour ajuster dynamiquement la pertinence des caractéristiques en utilisant l'Auto-Épi-Attention. Le jeu de données contient 569 échantillons, chacun décrit par 30 caractéristiques numériques extraites d'images de tissus mammaires. La tâche consiste à classer chaque échantillon comme bénin ou malin.

6.1.1 Architecture du modèle $epi - model_{BC}$

L'architecture comprend une couche d'entrée, une couche d'Auto-Épi-Attention pour moduler dynamiquement les caractéristiques selon les corrélations internes, et une couche de sortie avec une activation sigmoïde pour prédire la probabilité qu'une tumeur soit bénigne ou maligne.

Cette architecture est spécifiquement adaptée pour exploiter la puissance de l'Auto-Épi-Attention dans le diagnostic médical, permettant au modèle de se concentrer sur les caractéristiques les plus pertinentes pour chaque classe, améliorant ainsi l'explicabilité du modèle tout en maintenant les performances de classification.

6.1.2 Analyse de la pertinence des caractéristiques dans le Wisconsin Breast Cancer Dataset en utilisant l'Auto-Épi-Attention

L'Auto-Épi-Attention génère des vecteurs de pertinence des caractéristiques $a_i = [a_i^j]_{j=1}^{30} \in \mathbb{R}^{30}$ pour chaque échantillon, comme illustré à la figure 2 (gauche). Le jeu de données est structuré avec des cellules malignes dans les 212 premières lignes et des cellules bénignes dans les 357 suivantes. L'Auto-Épi-Attention permet de capturer les corrélations locales spécifiques à chaque classe, modulant dynamiquement la pertinence des caractéristiques.

Épi-Attention : Attention Dynamique Sensible au Contexte

- Rôle de la texture dans le diagnostic malin : La caractéristique `mean_texture` reçoit une attention élevée pour les tumeurs malignes, en accord avec son rôle dans le diagnostic du cancer Chitalia et Kontos (2019).
- Caractéristiques de surface dans les cellules bénignes : Les propriétés de surface, telles que `mean_fractal_dimension`, `mean_smoothness` et `mean_symmetry`, sont plus importantes dans les diagnostics bénins Wohl et al. (2023).

Ainsi, l'Épi-Attention capture la variabilité de l'importance des caractéristiques selon les classes, ajustant dynamiquement la pertinence des caractéristiques en fonction du contexte. De plus, elle permet d'éviter une pondération disproportionnée des caractéristiques pour les classes majoritaires, un problème commun aux modèles globaux Ünal et al. (2024).

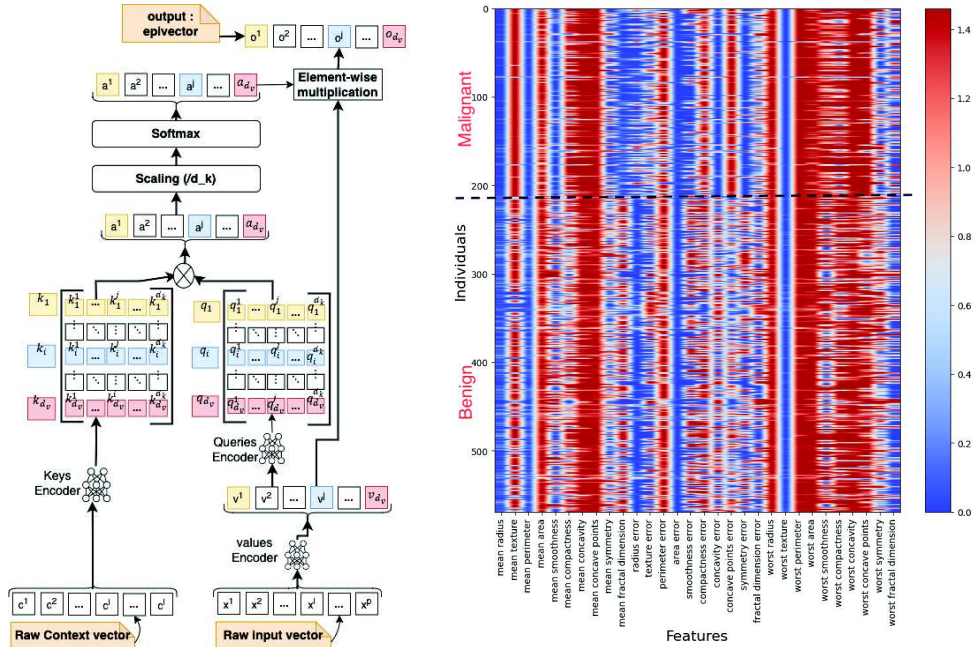


FIG. 2 – (gauche) Le mécanisme d'Épi-Attention, (droite) Importance des caractéristiques pour les 569 patients du Breast Cancer Dataset obtenue à partir de la couche d'auto-épi-attention de $epi - model_{CB}$.

6.2 Application 2 : Données Mixtes en Télémarcheting avec Bank Marketing Dataset

Le Bank Marketing Dataset est utilisé pour prédire la souscription des clients à un dépôt à terme, avec des caractéristiques catégorielles et numériques, idéal pour gérer les données mixtes. Le dataset est très déséquilibré (11% de succès), ce qui complique la capture des modèles pour la classe minoritaire. Nous avons développé l'épi-modèle ($epi - model_m$ voir figure 1) pour cette tâche de classification.

6.2.1 Architecture du modèle $epi - model_m$

L'architecture comprend deux couches d'entrée, l'une dédiée aux caractéristiques catégorielles encodées en one-hot et l'autre aux caractéristiques numériques. Une couche d'Épi-Attention ajuste ensuite dynamiquement la pertinence des caractéristiques numériques, permettant au modèle de moduler leur importance en fonction du contexte fourni par les données catégorielles. Enfin, une couche de sortie prédit si le client souscrira à un dépôt à terme, en s'appuyant sur les caractéristiques re-pondérées par la couche d'Épi-Attention.

La figure 3 (droite) (gauche) illustre l'importance des caractéristiques pour les attributs numériques dans le Bank Marketing Dataset, dérivée de la couche d'Épi-Attention de $epi - model_m$ (avec 91% de précision). En utilisant l'Épi-Attention, le modèle est capable d'adapter dynamiquement l'importance des caractéristiques pour les classes majoritaires et minoritaires, gérant plus efficacement le déséquilibre du jeu de données que les modèles traditionnels. Cela conduit à une meilleure explicabilité et garantit que le modèle capture les schémas critiques chez tous les clients, en particulier dans la classe sous-représentée avec plusieurs schémas hétérogènes.

6.3 Application 3 : Encodage Contextuel en Neuroimagerie avec ABIDE-II Dataset

Le jeu de données Autism Brain Imaging Data Exchange II (ABIDE-II) est un jeu de données de neuroimagerie à grande échelle conçu pour étudier les troubles du spectre de l'autisme (TSA) dans différentes populations et sites. Le jeu de données contient des données IRMf et IRM structurelles recueillies dans 20 sites internationaux, offrant une opportunité unique d'explorer la variabilité dans les données de neuroimagerie due aux différences de protocoles de numérisation et de distributions démographiques. Pour cette application, nous avons implémenté $epi - model_{ABIDE}$ dans la figure 1 pour encoder les séries temporelles (IRMf) tout en incorporant le contexte du site d'origine en tant qu'entrée supplémentaire. Ce modèle a été entraîné en utilisant une architecture siamoise pour capturer des encodages significatifs spécifiques aux sites pour la comparaison des séries temporelles.

6.3.1 Architecture du modèle $epi - model_{ABIDE}$

L'architecture comprend une couche d'entrée qui traite une série temporelle de 316 points représentant les données IRMf, ainsi qu'une couche supplémentaire qui reçoit un vecteur one-hot de dimension 20 pour indiquer le site de collecte. Une couche d'Épi-Attention ajuste ensuite dynamiquement la pondération des caractéristiques temporelles en fonction du contexte du site, permettant ainsi au modèle de moduler la pertinence des données en fonction des spécificités de chaque site.

Les encodages des séries temporelles obtenus sont ensuite comparés dans un cadre siamois, avec pour objectif d'apprendre des représentations robustes qui tiennent compte de la variabilité spécifique aux sites.

6.3.2 Focus sur les encodages des sites :

Alors que les applications précédentes se concentraient sur l'attention portée à des éléments spécifiques, cette application met en évidence les encodages générés pour les 20 sites sources. La figure 3 (droite) montre le premier plan principal d'une ACP appliquée aux encodages appris, révélant des clusters qui reflètent naturellement les similitudes entre les protocoles suivis par les différents sites. Le regroupement montre comment l'Épi-Attention capture les relations intrinsèques entre les sites, en fonction des protocoles de neuroimagerie et des variations démographiques. Ces résultats sont en accord avec les caractéristiques connues des sites Nielsen et al. (2013); Kunda et al. (2020).

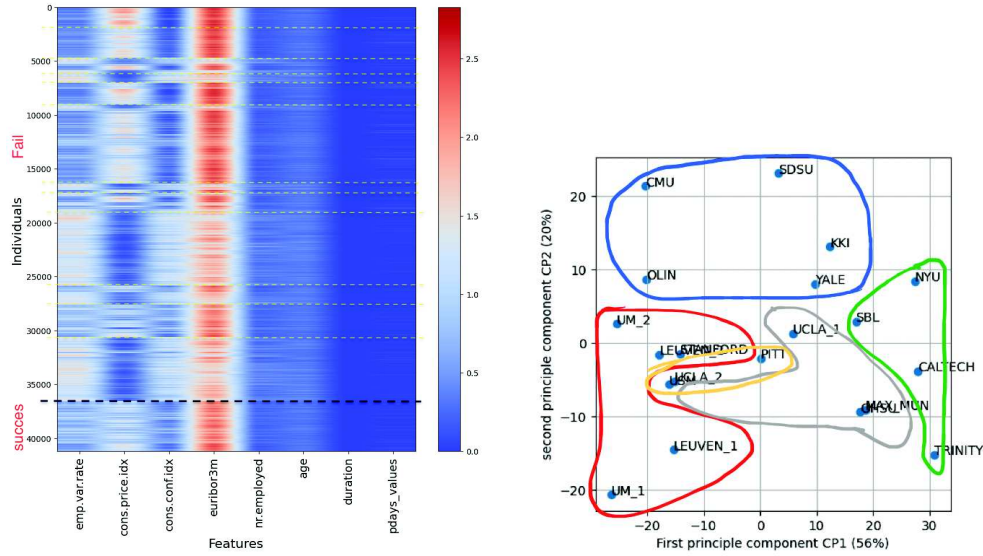


FIG. 3 – (gauche) Importance des caractéristiques pour les attributs numériques du Bank Marketing Dataset déséquilibré obtenue à partir de l'épi-attention de $\text{epi} - \text{model}_m$, (droite) Premier plan principal d'une ACP sur les encodages des sites avec net-K du $\text{epi} - \text{model}_{ABIDE}$.

7 Conclusion

Dans cet article, nous avons introduit l'Épi-Attention, un nouveau mécanisme qui ajuste dynamiquement la pertinence des caractéristiques en exploitant les corrélations internes et le contexte externe. Contrairement aux approches classiques à poids statiques, notre méthode permet une modulation fine et contextualisée des caractéristiques, ce qui améliore à la fois l'explicabilité des modèles et leur capacité à s'adapter à des scénarios complexes.

Nos expérimentations, couvrant des domaines variés tels que la santé, le télémarketing, et la neuroimagerie, ont démontré l'efficacité de l'Épi-Attention dans la capture des corrélations spécifiques aux classes, la gestion des données hétérogènes, et l'incorporation des variations

contextuelles. Ces résultats mettent en évidence son potentiel pour relever des défis réels où l’explicabilité des résultats et la précision des modèles sont essentielles.

En conclusion, l’Épi-Attention constitue une avancée prometteuse pour concevoir des modèles plus explicables et robustes. Les travaux futurs s’orientent vers son application à des ensembles de données encore plus variés (ex. images, données multimodales, ...) et vers l’intégration avec d’autres architectures modernes pour maximiser son potentiel.

Références

- Bo, S., Y. Zhang, J. Huang, S. Liu, Z. Chen, et Z. Li (2024). Attention mechanism and context modeling system for text mining machine translation. In *2024 6th International Conference on Data-driven Optimization of Complex Systems (DOCS)*, pp. 857–863. IEEE.
- Britz, D., M. Guan, et M. Luong (2017). Efficient attention using a fixed-size memory representation. In *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pp. 392–400.
- Chitalia, R. D. et D. Kontos (2019). Role of texture analysis in breast mri as a cancer biomarker : A review. *Journal of Magnetic Resonance Imaging* 49(4), 927–938.
- de Kok, J. W., F. van Rosmalen, J. Koeze, F. Keus, S. M. van Kuijk, J. Castela Forte, R. M. Schnabel, R. G. Driessen, T. T. van Herpt, J.-W. E. Sels, et al. (2024). Deep embedded clustering generalisability and adaptation for integrating mixed datatypes : two critical care cohorts. *Scientific Reports* 14(1), 1045.
- Hassanin, M., S. Anwar, I. Radwan, F. S. Khan, et A. Mian (2024). Visual attention methods in deep learning : An in-depth survey. *Information Fusion* 108, 102417.
- Kunda, M., S. Zhou, G. Gong, et H. Lu (2020). Improving multi-site autism classification based on site-dependence minimisation and second-order functional connectivity. Preprint, available at <https://doi.org/10.1101/2020.02.20.957670>.
- Lee, J., S. Kim, S. Kim, J. Park, et K. Sohn (2019). Context-aware emotion recognition networks. In *Proceedings of the IEEE/CVF international conference on computer vision*, pp. 10143–10152.
- Liu, D., H. Xu, J. Wang, Y. Lu, J. Kong, et M. Qi (2021). Adaptive attention memory graph convolutional networks for skeleton-based action recognition. *Sensors* 21(20), 6761.
- Liu, J., G. Wang, L.-Y. Duan, K. Abdiyeva, et A. C. Kot (2018). Skeleton-based human action recognition with global context-aware attention lstm networks. *IEEE Transactions on Image Processing* 27(4), 1586–1599, doi: 10.1109/TIP.2017.2785279.
- Liu, J., G. Wang, P. Hu, L.-Y. Duan, et A. C. Kot (2017). Global context-aware attention lstm networks for 3d action recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 1647–1656.
- Nielsen, J. A., B. A. Zielinski, P. T. Fletcher, A. L. Alexander, N. Lange, E. D. Bigler, J. E. Lainhart, et J. S. Anderson (2013). Multisite functional connectivity mri classification of autism : Abide results. *Frontiers in human neuroscience* 7, 599.
- Sahoo, S. et S. Chakraborty (2020). Learning representation for mixed data types with a non-linear deep encoder-decoder framework. arXiv preprint, available at <https://arxiv.org/abs/2008.00000>.

- [org/abs/2009.09634](https://arxiv.org/abs/2009.09634).
- Seo, M. J., A. Kembhavi, A. Farhadi, et H. Hajishirzi (2017). Bidirectional attention flow for machine comprehension. In *5th International Conference on Learning Representations, ICLR 2017, Toulon, France, April 24-26, 2017, Conference Track Proceedings*. OpenReview.net.
- Vaswani, A., N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, L. Kaiser, et I. Polosukhin (2017). Attention is all you need.
- Wohl, I., J. Sajman, et E. Sherman (2023). Cell surface vibrations distinguish malignant from benign cells. *Cells* 12(14), 1901, doi: 10.3390/cells12141901.
- Xiao, L., M. Li, Y. Feng, M. Wang, Z. Zhu, et Z. Chen (2024). Exploration of attention mechanism-enhanced deep learning models in the mining of medical textual data. arXiv preprint, available at <https://arxiv.org/abs/2406.00016>.
- Xing, C., W. Wu, Y. Wu, J. Liu, Y. Huang, M. Zhou, et W.-Y. Ma (2017). Topic aware neural response generation. In *Proceedings of the AAAI conference on artificial intelligence*, Volume 31.
- Yu, R., Y. Li, W. Lu, et L. Cao (2022). Tri-attention : Explicit context-aware attention mechanism for natural language processing. arXiv preprint, available at <https://arxiv.org/abs/2211.02899>.
- Zhong, P., D. Wang, et C. Miao (2019). Knowledge-enriched transformer for emotion detection in textual conversations. In K. Inui, J. Jiang, V. Ng, et X. Wan (Eds.), *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, Hong Kong, China, pp. 165–176. Association for Computational Linguistics, doi: 10.18653/v1/D19-1016.
- Ünal, S., O. Günay, I. Akkurt, K. Gunoglu, et H. Tekin (2024). A comparative study on breast cancer classification with stratified shuffle split and k-fold cross validation via ensembled machine learning. *Journal of Radiation Research and Applied Sciences* 17(4), 101080, doi: <https://doi.org/10.1016/j.jrras.2024.101080>.

Summary

This paper presents Epi-Attention, a context-aware mechanism that dynamically adjusts feature relevance in neural networks to capture local correlations and class-specific patterns. Additionally, Self-Epi-Attention is proposed to enhance internal feature consistency. Applications in medical diagnosis, marketing, and neuroimaging demonstrate improved model interpretability and adaptability without compromising performance.