



Projet d'accord-cadre sur le déploiement et l'utilisation de l'intelligence artificielle (IA) dans la Fonction publique – version 4

15 septembre 2026

Déclaration liminaire

Depuis le début de cette négociation, la CGT dénonce un calendrier beaucoup trop contraint.

Quelques réunions, avec l'été au milieu, pour négocier un accord qui doit couvrir les trois versants de la Fonction publique et traiter de l'emploi, du travail, de la santé, des qualifications, des données, des usager·ères, de l'environnement et des responsabilités : ce n'est pas sérieux.

D'autant que pendant que nous négocions, les déploiements, eux, accélèrent en laissant sur le côté de la route les négociations avec les organisations syndicales.

L'exemple le plus symptomatique est l'annonce de la généralisation de l'assistant IA à l'ensemble des agent·es de l'État, après une expérimentation auprès de 10 000 agent·es dont les conclusions ne sont pas que positives et ce, deux jours avant le début de l'ouverture de cet accord et alors même que les agent·es n'ont pas bénéficié d'une formation préalable.

Et cette course au déploiement intervient dans un contexte particulier.

Au niveau mondial, les principaux acteurs et principales actrices de l'IA s'interrogent désormais eux-mêmes et elles-mêmes sur la nécessité de ralentir le développement des modèles les plus avancés, le temps de progresser sur leur alignement, leur évaluation et leur contrôle.

Ce n'est pas le meilleur moment pour généraliser et accélérer un déploiement avant même qu'un accord national soit signé.

Pour la CGT, une question doit donc précéder toutes les autres : pourquoi introduire de l'IA ?

Pour quelle tâche ? Pour répondre à quel problème identifié dans le travail réel ? Et pourquoi une IA plutôt qu'une réponse humaine, organisationnelle ou technique ?

C'est le rôle de l'étude d'opportunité que nous demandons avant toute décision. Elle doit partir du travail réel : des tâches effectivement réalisées, des difficultés rencontrées, des besoins des agent·es et des usager·ères. Elle doit examiner les différentes réponses possibles, humaines, organisationnelles ou techniques, avant de présupposer que l'IA est la solution.

Et si l'opportunité est démontrée, une étude d'impact contradictoire doit précéder tout déploiement. Elle doit mesurer les transformations du travail réel : tâches supprimées, créées ou déplacées, charge et

intensification du travail, autonomie, collectifs, emploi et qualifications, santé, mais aussi effets sur les usager·ères, les données et l'environnement. Les risques pour la santé et les conditions de travail doivent être intégrés au DUERP et donner lieu à des mesures de prévention dans les PAPRIPACT.

Ses résultats doivent être présentés aux instances compétentes et pouvoir conduire à modifier, suspendre ou abandonner le projet.

Les enjeux sur l'emploi imposent cette exigence.

Le rapport IGF-IGA-IGAS d'avril 2026, Le déploiement de l'intelligence artificielle dans les administrations publiques : analyse comparée, enjeux et conditions de réussite, estime que 20 % des agent·es public·ques, soit 1,13 million de personnes, occupent des emplois exposés à l'IA. Plus de 700 000 agent·es sont dans les catégories les plus ou significativement exposées

Ces chiffres interdisent précisément de lever les inquiétudes légitimes des agent·es sur l'emploi et dont nous n'avons aucune garantie dans le projet d'accord.

On nous parle d'un gain de temps sans réellement le démontrer, et pendant ce temps, d'anciens membres du gouvernement, candidats à la présidentielle, annoncent qu'ils supprimeront des centaines de milliers de fonctionnaires grâce à l'IA, entre autres.

Dans son rapport L'intelligence artificielle dans les politiques publiques : l'exemple du ministère de l'Économie et des Finances, publié en 2024, la Cour des comptes étudie treize systèmes d'IA en service. Le repositionnement vers des tâches à plus forte valeur ajoutée n'est précisément documenté que dans un seul cas, concernant vingt-cinq agent·es.

Elle relève en revanche de nouvelles tâches de préparation, de tri, de relecture et de contrôle.

Nous connaissons cette histoire.

Chaque vague de numérisation devait libérer du temps. Les agent·es ont surtout connu l'accélération des rythmes, l'augmentation des objectifs, les réductions d'effectifs et l'intensification du travail.

L'IA ne doit pas devenir l'étape suivante.

Un gain de temps n'est pas un gain pour l'agent·e s'il sert à lui demander davantage de dossiers ou à supprimer le poste d'un·e collègue d'à côté.

Nous exigeons donc qu'aucun gain lié à l'IA ne puisse justifier suppressions ou non-remplacements de postes, réductions d'effectifs, redéploiements contraints, déqualifications ou augmentation des normes et volumes de production.

Même exigence sur la responsabilité.

Nous refusons tout transfert de responsabilité vers les agent·es.

L'administration choisit l'outil, le fournisseur, les conditions de déploiement et l'organisation du travail. Elle doit en assumer la responsabilité.

« L'humain aux commandes » ne peut pas signifier l'agent·e responsable au bout d'une chaîne qu'il ou elle ne maîtrise pas.

Un contrôle réel exige du temps, de la formation, des compétences, de la traçabilité et l'accès aux sources.

Et la formation doit permettre de comprendre les limites et les risques de l'outil, pas seulement d'apprendre à s'en servir.

Et les conséquences d'une erreur ne sont pas théoriques. Dans la Fonction publique hospitalière, une erreur dans une information médicale ou une aide au diagnostic peut avoir des conséquences extrêmement graves.

Un agent·e ne peut donc être rendu·e responsable des erreurs, biais, hallucinations ou défaillances qu'il ou elle ne pouvait raisonnablement identifier ou contrôler.

Un système IA peut fonctionner comme prévu tout en produisant des discriminations. À France Travail, si l'algorithme cible davantage certaines catégories, notamment les femmes précaires, il organise leur surcontrôle. Le problème n'est alors ni l'erreur de l'IA ni celle de l'agent·e : il est dans les choix faits en amont. Corriger le modèle ne suffit pas si la finalité reste la même.

Et ces choix ne sont pas ceux de l'agent·e et ne font l'objet d'aucun dialogue social.

C'est précisément pour cela que nous avons besoin de garanties collectives et opposables.

Nous ne voulons pas d'une multiplication de chartes ou de guides unilatéraux. Nous voulons des accords collectifs négociés et signés, dans les ministères, les collectivités et les établissements, notamment hospitaliers.

Des accords qui garantissent l'emploi, les qualifications, la santé, la formation, les données et les droits des usager·ères. La surveillance, la notation, le profilage et l'évaluation automatisée des agent·es doivent être interdits. Le droit d'alerte, la suspension et la réversibilité doivent être garantis.

Une charte ne remplace pas un accord collectif.

La Fonction publique n'a pas besoin d'un accord qui accompagne l'accélération des déploiements.

Elle a besoin d'un accord qui donne aux agent·es et à leurs représentant·es des droits pour intervenir avant la décision, contrôler les effets sur le travail et arrêter un système lorsque c'est nécessaire.

Nous refusons que l'IA soit considérée comme une évolution inéluctable à laquelle les agent·es devraient simplement s'adapter.

Le dialogue social doit pouvoir peser sur les choix technologiques eux-mêmes.

Y compris pour décider de ne pas déployer.