

Büyük Dil Modelleri Çağında Türkiye’de Felsefe Eğitimi: Nöropedagojik Bir Sıçrama Fırsatı ve Epistemolojik Bir Hesaplaşma

Öz

Büyük dil modellerinin yükseköğretime girişi çoğu üniversitede tek bir soruya indirgenmiştir: Öğrenci ödevini kendisi mi yazdı, yoksa yapay zekâya mı yazdırdı? Bu soru önemsiz değildir; fakat daha derindeki krizi gizlemektedir. Türkiye’de felsefe eğitimi, büyük dil modellerinden önce de tek yönlü ders anlatımı, eskimiş izlenceler, güncel uluslararası literatürle zayıf temas, betimleyici tezler, düzensiz danışmanlık ve yazılı ürün merkezli değerlendirme gibi sorunlarla karşı karşıyaydı. Büyük dil modelleri bu sorunları yaratmadı; onların artık görmezden gelinemeyecek kadar belirginleşmesini sağladı. Bu makale, büyük dil modellerinin Türkiye’de felsefe öğretimi ve araştırması için yalnızca akademik dürüstlük riski değil, tarihsel bir telafi ve sıçrama imkanı sunduğunu savunmaktadır. Oakley, Rogowsky ve Sejnowski’nin nöropedagojik çerçevesinden hareketle öğrenmenin “anladığını hissetmek” değil, bilgiyi uzun süreli bellekte örgütlemek, geri çağırma ve yeni durumlara aktarabilmek olduğu ileri sürülmektedir. Buna göre büyük dil modelleri hazır cevap makineleri olarak değil; kişiselleştirilmiş geri çağırma, aralıklı tekrar, argüman provası, kaynak doğrulama ve sözlü savunma araçları olarak kullanılmalıdır. İzlencelerden ders kitaplarına, lisansüstü öğrenci kabulünden doktora yeterliğine ve Tez İzleme Komitelerine kadar bütün yapı bu yeni nöropedagojik ve teknolojik gerçekliğe göre yeniden tasarlanmalıdır.

Anahtar Kelimeler: Büyük dil modelleri, felsefe eğitimi, nöropedagoji, bellek, epistemik özerklik

Philosophy Education in Türkiye in the Age of Large Language Models: A Neuropedagogical Opportunity for a Leap Forward and an Epistemological Reckoning

Abstract

The arrival of large language models in higher education has often been reduced to a single question: Did the student write the assignment, or was it generated by artificial intelligence? Although important, this question conceals a deeper crisis. Philosophy education in Türkiye had already been marked by one-way lecturing, outdated syllabi, weak engagement with current international scholarship, descriptive theses, irregular supervision, and assessment practices centred on take-home written products. Large language models did not create these problems; they made them impossible to ignore. This article argues that such models offer not only risks to academic integrity but also a historically significant opportunity to compensate for structural disadvantages and enable a leap forward in philosophy education and research. Drawing on the neuropedagogical framework developed by Oakley, Rogowsky, and Sejnowski, it maintains that learning is not merely the feeling of understanding but the organisation, retrieval, and transfer of knowledge stored in long-term memory. Large language models should therefore be used not as answer-generating machines but as tools for personalised retrieval practice, spaced repetition, argumentative rehearsal, source verification, and oral defence. Syllabi, course materials, assessment, graduate admissions, doctoral qualification examinations, thesis monitoring committees, and supervision should all be redesigned in accordance with this new technological and neuropedagogical reality.

Keywords: Large language models, philosophy education, neuropedagogy, memory, epistemic autonomy

Giriş

Büyük dil modelleri üniversiteye girdiğinden beri yükseköğretim dünyasında tuhaf bir telaş vardır. Öğrenciler yapay zekâyla ödev yazmakta, öğretim üyeleri metnin kime ait olduğunu sezgileriyle anlamaya çalışmakta, üniversiteler etik yönergeler yayımlamaktadır. İnsan kurumlarının teknolojik bir kırılmayla karşılaştığında önce paniğe kapılıp ardından bir belge hazırlaması şaşırtıcı değildir.¹

Fakat tartışmanın başlangıç noktası yanlıştır. Sanki üniversite eğitimi büyük dil modellerinden önce kusursuz çalışıyormuş da bu sistemler gelip düzeni bozmuş gibi davranılmaktadır. Oysa büyük dil modelleri üniversiteyi bozmadı; zaten sorunlu olan yerleri görünür kıldı. “Hoca anlatır, öğrenci dinler; öğrenci eve gider, uzun bir yazılı ödev hazırlar; bu metin de öğrencinin ne öğrendiğini gösterir” modeli çevrim içi video dersler, açık ders kaynakları ve sayısal kütüphanelerle çoktan aşınmıştı.²

Makine zekâsı ile eğitim arasındaki bağ da yeni değildir. Turing’in 1948 tarihli Intelligent Machinery raporu, yetişkin zihnini doğrudan kopyalamak yerine eğitilebilen bir çocuk-makine tasarlama düşüncesini tartışıyordu (Turing, 1948). Cahit Arf ise 1959’da “anlamadan belleme” alışkanlığını eleştiriyor, kitapların ve başka insanların zihnin yardımcı hafızası gibi çalışabileceğini belirtiyor ve zekânın ayırt edici yönünü yeni durumlara uyum sağlama kapasitesinde buluyordu (Arf, 1959). Yetmiş yıl sonra öğrenciden hala büyük ölçüde hazır bilgiyi yeniden üretmesini beklememiz, teknolojik yetersizlikten çok kurumsal yenilenmenin gecikmiş olmasıyla ilgilidir.

Türkiye’deki yapay zekâ felsefesi literatürü de artık yalnızca “makinelere düşünebilir mi?” sorusuyla sınırlı değildir. Akademik yazarlık, anlama, semantik içerik, bilinç, doğal zekânın dönüşümü ve üretken yapay zekânın kavramsal statüsü ayrı ayrı ele alınmaktadır (Güven ve Çoban Sarı, 2026; Erciyes, 2024). Bu ayrımlar eğitim tartışması için de zorunludur. Bir sistemin öznel deneyime sahip olup olmadığı, onun öğrenciye Hume’un nedensellik eleştirisini daha doğru kurdurup kuramayacağıyla aynı soru değildir. Metafizik tartışmaların gerekli pedagojik yenilikleri ertelemenin gerekçesine dönüşmesine izin vermemek gerekir.³

Büyük Dil Modellerinin Epistemik Statüsü

Büyük dil modelleri hakkında iki kolaycı anlatı vardır. Birincisi, bu sistemlerin insan benzeri anlayışa ve muhakemeye ulaştığıdır. İkincisi, yalnızca bir sonraki tokeni tahmin ettikleri için ciddi hiçbir bilişsel kapasiteye sahip olamayacaklarıdır. İkisi de slogan olarak kullanışlı, açıklama olarak yetersizdir. Bu makalede “büyük dil modeli” terimi, yalnızca temel modelleri değil, onların çevresinde kurulan araç bağlantılı uygulama katmanını da kapsayacak biçimde kullanılmaktadır; eğitimsel kullanım iddialarının çoğu fiilen bu birleşik sistemlere ilişkindir.

Millière ve Buckner’a göre bir sistemin eğitim hedefini onun bütün işlevlerinin eksiksiz betimi saymak bir “yeniden betimleme yanılgısı”dır. Bir modelin bir sonraki tokeni tahmin etmek üzere eğitilmiş olması, eğitimden sonra yaptığı her şeyin yalnızca token tahmini diye açıklanabileceği anlamına gelmez (Millière ve Buckner, 2026). Buna karşılık yüksek başarıyı doğrudan anlama, inanç veya bilinç kanıtı saymak da acelecidir. Bir öğrencinin Rorty hakkında düzgün bir paragraf yazması, Rorty’yi anladığını tek başına kanıtlamıyorsa, bir modelin aynı paragrafı üretmesi de zihinsel hayatı hakkında fazla şey göstermez.

Güven, anlamanın bağlama duyarlılık, çıkarımsal örgütlenme, semantik temellendirme, normatif katılım ve özne merkezli farkındalık gibi farklı katmanlara ayrılması gerektiğini savunur. Güncel modeller ilk katmanların bazılarında güçlü olabilirken daha kuvvetli anlama ölçütlerini karşılamıyor olabilir (Güven, 2026). Sarı’nın “nedensel risk” kavramı bu ayrımı keskinleştirir: Bir temsilin doğru veya yanlış olmasının sistemin kendi işleyişi

¹ “AI”, İngilizce artificial intelligence ifadesinin kısaltmasıdır; bu makalede “yapay zekâ” ve gerektiğinde “YZ” kullanılmaktadır. Yapay zekânın üzerinde uzlaşmış tek bir tanımı yoktur; en geniş anlamıyla, geçmişte insan zekâsı gerektirdiği düşünülen görevleri (algılama, öğrenme, dil işleme, planlama, problem çözme) yerine getiren hesaplamalı sistemleri kapsar. Kural tabanlı sistemlerden derin öğrenmeye ve üretken sistemlere uzanan bu alanda her sistem öğrenmez veya dil kullanmaz; bu nedenle “yapay zekâ” ile “büyük dil modeli”ni eş anlamlı kullanmak teknik bakımdan yanlıştır (Lake vd., 2017; Monett vd., 2020).

² “LLM”, İngilizce large language model ifadesinin kısaltmasıdır; Türkçede “büyük dil modeli” veya “BDM”. Büyük dil modelleri, çok büyük metin (ve çoğu durumda kod) derlemleri üzerinde eğitilen, dilsel dizilerde bir sonraki tokenin olasılığını hesaplayan yapay sinir ağlarıdır ve çoğunlukla dikkat düzenine dayalı dönüştürücü mimarisini kullanır. Ön eğitimden sonra yönerge izleme, diyalog ve güvenlik amaçlarıyla ayrıca eğitilirler; ayırt edici özellikleri, çok çeşitli görevleri tek bir arayüzden yürütebilmeleri ve yeni görevlere az örnekle uyarlanabilmeleridir. Bu makalede terim, yalnızca temel modelleri değil, ChatGPT, Claude, Gemini ve Grok gibi çok kipli, araç kullanabilen birleşik uygulamaları da kapsayacak biçimde kullanılmaktadır; çünkü gündelik hayatta karşılaşılan nesne çoğunlukla bu birleşik sistemlerdir (Millière ve Buckner, 2026; Vaswani vd., 2017).

³ Yapay zekânın gerçekten düşünüp düşünemeyeceği, bilince sahip olup olamayacağı, insanlığı tehdit edip etmeyeceği ve emeği ne ölçüde ikame edeceği önemli fakat birbirinden farklı soru kümeleridir; sırasıyla zihin felsefesi ve metafizik, yapay zekâ güvenliği ile emek iktisadı ve teknoloji politikasının konularıdır. Bunları tek bir “yapay zekâ sorunu” altında toplamak kavramsal bulanıklık yaratır. Üstelik hiçbirinin cevabı, büyük dil modellerinin geri çağırma, açıklama veya argüman çözümlemesi bakımından eğitimsel olarak yararlı olup olmadığı tek başına belirlenmez. Bu makale söz konusu tartışmaları reddetmemekte, yalnızca yöntemsel bir kapsam sınırlamasıyla askıya almaktadır.

bakımından gerçek bir bedeli yoksa, akıcı çıktı tek başına içsel semantik içerik kanıtı sayılamaz (Sarı, 2026). Model yanlış bir Quine yorumu verdiğinde epistemik bedeli model değil, öğrenci veya araştırmacı öder. Bu nedenle doğrulama sorumluluğu modele devredilemez.

Hicks, Humphries ve Slater, büyük dil modellerinin yanlışlarını “halüsinasyon” olarak değil, Frankfurtçu anlamda hakikate kayıtsız metin üretimi olarak değerlendirmektedir (Hicks vd., 2024). Bu yorum bütünüyle kabul edilmese bile modelin akıcı anlatımıyla doğruluğunu ayırmak zorunludur. Çözüm yasaklamak değil, doğrulamayı öğretmektir.

Lake ve arkadaşlarının insan benzeri öğrenme çözümlemesi de performans ile öğrenme biçimini ayırmayı gerektirir. İnsanlar az örnekten, ön bilgiye dayanarak, nedensel modeller kurarak ve yeni amaçlara hızla uyum sağlayarak öğrenir; derin öğrenme sistemleri ise birçok görevde çok daha fazla veriye ve eğitime ihtiyaç duyar (Lake vd., 2017). Bu fark büyük dil modellerinin eğitimsel değerini ortadan kaldırmaz. Model öğrencinin yerine epistemik özne değildir; öğrencinin daha iyi soru sormasına, karşı örnek üretmesine ve itirazlara hazırlanmasına yardım eden bir araçtır.

Bu bölümün pedagojik argümana taşıdığı yük, açıkça adlandırılması gereken tek bir köprü öncüle indirgenebilir: Bir sistemin eğitimsel yararlılığı, o sistemin güçlü anlamda anlamasını, inanç taşımasını veya bilinçli olmasını gerektirmez; yararlılık, sistemin öğrencinin geri çağırma, ayırtırma, doğrulama ve savunma pratiklerini güvenilir biçimde tetikleyip tetiklemediğiyle ölçülür. Makalenin izleyen bölümleri yalnızca bu zayıf öncüle dayanır; daha güçlü metafizik iddiaların doğruluğu veya yanlışlığı buradaki önerileri etkilemez.

Nöropedagojik Zemin: Bellek, Geri Çağırma ve Öğrenci Farklılıkları

Nöroeğitim, öğrenme ve öğretim tasarımı sinirbilim, bilişsel psikoloji ve eğitim araştırmalarının bulgularıyla ilişkilendiren disiplinler arası alandır. Buradaki amaç beyin görüntülerinden aceleci reçeteler çıkarmak değil; bellek, dikkat, bilişsel yük, geri çağırma ve otomatikleşme hakkındaki bulguları pedagojik kanıtlarla birlikte kullanmaktır.

Oakley, Rogowsky ve Sejnowski'nin *Uncommon Sense Teaching: Practical Insights in Brain Science to Help Students Learn* (Sağduyuya Aykırı Öğretim: Öğrencilerin Öğrenmesine Yardımcı Olmak İçin Beyin Biliminden Pratik İçgörüler) adlı kitabı bu bakımdan temel bir çerçeve sunar. Terrence J. Sejnowski'nin hesaplamalı sinirbilimin ve yapay sinir ağları araştırmalarının öncü isimlerinden biri olması, insan öğrenmesi ile makine öğrenmesi arasındaki bağın yüzeysel benzetmelerle değil, iki alanı da bilen bir araştırmacının katkısıyla kurulmasını sağlar.

Kitabın ana mesajı basittir: Öğrencinin bir açıklamayı takip etmesi, öğrenmiş olması değildir. Öğrenme, bilgilerin uzun süreli bellekte bağlantılar kurması, zaman içinde güçlenmesi ve uygun durumda geri çağırılabilmesidir (Oakley vd., 2021). Büyük dil modelleri burada iki zıt etki yaratabilir. Akıcı açıklamalar öğrencide sahte bir anlama hissi oluşturabilir; fakat doğru kullanıldıklarında geri çağırma, aralıklı tekrar, karıştırmalı alıştırma ve düzeltici geri bildirim için güçlü araçlara dönüşebilir. Bu ayrım yakın dönem deneysel kanıtlarla da uyumludur: Türkiye’de bir lisede yürütülen alan deneyinde korumasız GPT-4 arayüzü öğrencilerin yapay zekâ destekli alıştırma performansını artırmasına rağmen, erişim kaldırıldığında bağımsız sınav performansını düşürmüştü; öğrenmeyi korumak için tasarlanan GPT Tutor koşulunda ise bu olumsuz etki büyük ölçüde azaltılmıştır (Bastani vd., 2025).

Bu çerçevenin dayandığı bulgular popüler bir sentezden ibaret değildir. Bir öğrenme malzemesini yeniden çalışmak yerine sınanmanın, gecikmeli sınamada uzun süreli kalıcılığı belirgin biçimde artırdığı deneysel olarak gösterilmiştir; aynı çalışmada beş dakikalık kısa gecikmede tekrar çalışmanın öne geçtiği, üstünlüğün ancak günler sonraki sınamada ortaya çıktığı bulunmuştur (Roediger ve Karpicke, 2006). Zamana yayılmış çalışmanın toplu çalışmaya üstünlüğü ise 317 deneyi kapsayan geniş bir meta-analizle belgelenmiştir (Cepeda vd., 2006). Bu davranışsal bulguların sinirsel düzeydeki bir karşılığı da vardır: Bellek izleri, uyku sırasında ve özellikle yavaş dalga uykusunda yeniden etkinleşme yoluyla sistem düzeyinde pekiştirilir (Rasch ve Born, 2013). Bu mekanizma doğrudan aralıklı tekrarın nedeni olarak sunulamaz; çünkü aralık etkisi uyanık dönemlerdeki aralıklarda da görülür ve uykunun aralık etkisindeki nedensel payı hâlâ tartışmalıdır. Yine de bulgu, dağıtılmış çalışmanın ve dinlenme aralıklarının neden basit bir “anladığını hissetme” meselesi olmadığını gösterir: Kalıcılık, zaman içinde işleyen konsolidasyon süreçlerine bağlıdır.

Öğrenciye “Hume’un nedensellik anlayışını açıkla” diye hazır bir metin yazdırmak yerine şu çalışma yaptırılabilir: “Metne bakmadan Hume’un zorunlu bağlantı eleştirisini üç adımda kurmamı sağla; yanlışıma hemen söyleme, önce ipucu ver; üç gün sonra aynı soruyu Mill’in tümevarım anlayışıyla karıştırarak yeniden sor.” Benzer biçimde öğrenciye Kuhn’un paradigma kavramını özetletmek yerine olağan bilim, bunalım ve bilimsel devrim arasındaki ilişkiyi farklı örneklerde yeniden kurması istenebilir. Böyle kullanıldığında büyük dil modeli hazır bilgi dağıtan bir otomat değil, belleği çalıştıran kişisel antrenördür.

Dunlosky ve arkadaşlarının kapsamlı değerlendirmesi, geri çağırma sınamalarını ve zamana yayılmış çalışmayı yüksek faydalı teknikler arasında gösterirken, tekrar okuma ve altını çizme gibi yaygın alışkanlıkların daha sınırlı kaldığını ortaya koymuştur (Dunlosky vd., 2013). Öğrencinin Quine makalesinin yarısını fosforlu kalemle boyaması zihinsel bir olay değildir. Metne bakmadan çözümleyici-sentetik ayırımına yöneltilecek temel itirazları kurabiliyorsa öğrenme başlamıştır.

Karşı yöndeki kanıtlar da ciddiye alınmalıdır. Randomize bir deneyde ChatGPT desteği yazma görevindeki anlık performansı yükseltmiş, fakat içsel motivasyonu artırmamış ve öz-düzenleyici öğrenme süreçlerinde “üstbilişsel tembellik” olarak adlandırılan bir bağımlılık örüntüsüyle ilişkilendirilmiştir (Fan vd., 2025). Korelasyonel düzeyde ise tek bir kesitsel araştırmada (N = 666), yoğun yapay zekâ aracı kullanımının bilişsel yük devri üzerinden daha düşük eleştirel düşünme puanlarıyla ilişkili bulunduğu bildirilmiştir (Gerlich, 2025). Bu bulgular makalenin tezini çürütmez; sınırını çizer: Büyük dil modellerinin pedagojik değeri aracın kendisinde değil, kullanım tasarımıdadır. Doğrulama, geri çağırma ve savunma yükünü öğrencide bırakan tasarımlar bu riskleri azaltır; cevabı hazır sunan kullanım biçimleri onları büyütür.

Öğrenciler çalışma belleği kapasitesi, ön bilgi, akademik dil, kaygı ve öğrenme hızı bakımından farklıdır. Oakley ve arkadaşlarının “yarış arabası” ve “yürüyüşçü” benzetmesi, hızlı işleyen öğrenciler ile daha yavaş fakat daha geniş bağlantılar kuran öğrenciler arasındaki farkı görünür kılar. Bu benzetme, ampirik desteği bulunmayan “öğrenme stilleri” sınıflamalarına değil, çalışma belleği ve işleme hızındaki ölçülebilir farklara gönderme yapar. Tek hızda ilerleyen ders anlatımı hızlı öğrenciyi ölçü, geri kalanını sapma sayar. Büyük dil modeli ise aynı kavramı farklı güçlük düzeylerinde açıklayabilir, İngilizce metne Türkçe ön okuma rehberi hazırlayabilir, çekingen öğrenciye düşük riskli bir soru alanı açabilir ve eksik ön bilgiye göre alıştırma üretebilir. Bu yalnızca kişiselleştirme değil, epistemik adalet meselesidir.

Felsefe Bilgi Deposu Değil, Düşünme Zanaatıdır

Türkiye’de felsefe öğretiminin temel kusurlarından biri, felsefeyi düşünürlerin görüşleri hakkında bildirimsel bilgiye indirgemesidir. Öğrenci Hume’un izlenim-fikir ayırımını, William James’in bilinç akışını, Charles Sanders Peirce’in pragmatik ilkesini, Mill’in zarar ilkesini, Quine’in inanç ağı görüşünü, Kuhn’un paradigma kavramını, Dennett’in yönelimsel tutumunu veya Rorty’nin ayna benzetisi eleştirisini anlatabilir. Fakat bir argümanı öncül-sonuç biçiminde kuramaz, gizli varsayımı bulamaz, rakibinin görüşünü güçlendiremez ve kendi tezine ciddi bir itiraz üretemez. Bu durumda felsefe öğrenmemiş, felsefe hakkında veri depolamıştır. Bu teşhis, henüz sistematik verilerle belgelenmemiş fakat alan içinde yaygın biçimde paylaşılan bir mesleki gözleme dayanmaktadır; Yükseköğretim Kurulu Ulusal Tez Merkezi üzerinden yapılacak basit bir başlık taraması bile betimleyici tez kalıbının yaygınlığını sınamak için elverişli bir başlangıç noktasıdır.

Felsefi yeterlilik aynı zamanda işlemseldir. Kavram ayırımı yapmak, karşı örnek kurmak, metinsel kanıt göstermek ve güçlü bir itiraza cevap vermek tekrar yoluyla öğrenilir. Öğrenci modele “Dennett’in yönelimsel tutumunu anlat” demekle yetinmemelidir. Modelden yönelimsel tutum lehine en güçlü savunuyu kurmasını, ardından araçsalcılık itirazını geliştirmesini, öğrencinin cevabındaki örtük öncülü bulmasını ve sözlü savunma provası yaptırmasını istemelidir. Peirce için kaçırma çıkarımı, tümevarım ve tümdengelim ayrıştırılabilir; Rorty için temsilcilik eleştirisinin gerçekçilik karşısındaki bedeli sorgulanabilir; William James için bilinç akışı ile alışkanlık arasındaki ilişki yeniden kurulabilir. Yakın tarihli ampirik bir çalışma da ChatGPT ile etkileşimin, özellikle öğrencinin etkin katılımı üzerinden, karmaşık eleştirel düşünme performansı ile ilişkili olduğunu göstermektedir (Suriano vd., 2025).

Kuşcu’nun uyarısı burada önemlidir: Tehlike yalnızca makinenin insanlaşması değil, insan düşüncesinin hız, ölçülebilirlik ve çıktı üretimine indirgenerek hesaplayıcı modele benzemesidir (Kuşcu, 2026). Felsefe eğitimi bu eğilime direnmelidir. Model öğrencinin değerlendirme yetisini ikame etmemeli; onu daha çok gerekçe sunmaya, karşı görüşü adil temsil etmeye ve kendi sonucunun bedelini düşünmeye zorlamalıdır.

Her ders anlatımı kötü değildir. Özellikle yeni ve zor konularda iyi yapılandırılmış kısa ve doğrudan açıklama değerlidir. Sorun, uzun, tek yönlü, geri çağırmasız ve uygulamasız ders anlatımıdır. Sınıf zamanı bilgi aktarımı için değil, felsefi işlem için kullanılmalıdır: Ders öncesinde temel bilgi ve ön okuma; sınıfta argüman çözümleme, karşılaştırma ve sözlü savunma; ders sonrasında geri çağırma ve aralıklı tekrar. Etkin öğrenmenin fen, teknoloji, mühendislik ve matematik alanlarında olduğu kadar beşerî ve sosyal bilimlerde de geleneksel ders anlatımından daha iyi sonuçlar verdiği ilişkin üst çözümlemeler bulunmaktadır (Freeman vd., 2014; Kozanitis ve Nenciovici, 2023).

İzlen, Ders Materyali ve Değerlendirme Reformu

İzlen bürokratik bir konu listesi değil, dersin epistemik mimarisi olmalıdır. Her hafta öğrencinin hangi problemi inceleyeceği, hangi birincil metni okuyacağı, hangi kavramı geri çağıracağı, hangi argümanı kuracağı, büyük dil modelini nerede kullanacağı ve hangi performansı sınıfta savunacağı açıkça belirtilmelidir.

Bilim felsefesi dersinde haftalar yalnızca “Quine”, “Kuhn” ve “Rorty” diye sıralanmamalıdır. Bunun yerine şu problemler kurulabilir: Gözlem kuramdan bağımsız olabilir mi? Bilimsel değişim yalnızca kanıt birikimiyle açıklanabilir mi? Farklı paradigmlar aynı dünyayı mı betimler? Hakikat ile gerekçelendirme arasında nasıl bir ilişki vardır? Öğrenci bu düşünceleri yalnızca anlatmamalı; aralarındaki uyumsuzlukları kurup savunabilmelidir.

Her izlenecede görev bazında açık yapay zekâ politikası bulunmalıdır. Sınıf içi yazma ve canlı sözlü sınav yapay zekâ kullanımına kapalı olabilir; ön okuma, dil desteği ve soru hazırlığı sınırlı kullanıma açılabilir; model cevabındaki hataları bulma, kaynak doğrulama ve argüman yeniden kurma görevlerinde ise yapay zekâ kullanımı zorunlu olabilir. “Kullanmayın” deyip öğrencinin kullandığını bilmek politika değil, kurumsal göz kapamadır.

Bu yaklaşım değerlendirmeyi yapay zekâyı kapalı ve yapay zekâyı açık etkinliklerin bilinçli bir bileşimi olarak kurmayı gerektirir. Birinci gruptaki görevler öğrencinin bağımsız geri çağırma, metin çözümleme ve sözlü muhakeme kapasitesini ölçerken, ikinci gruptakiler öğrencinin araçla çalışırken kaynakları doğrulama, model çıktısındaki boşlukları belirleme ve kendi kararının sorumluluğunu üstlenme becerisini sınamalıdır. Böylece amaç, öğrenciyi yapay zekâdan yalıtılmış hayalî bir dünyada sınamak ya da her işi modele bırakmak değildir. Amaç, bağımsız felsefi yeterlilik ile araç destekli araştırma yeterliliğini birbirinden ayırarak ikisini de geliştirmektir. Öğrencinin yalnızca sonuç metni değil, soruyu nasıl daralttığı, hangi kaynakları elelediği, modelin hangi önerilerini reddettiği ve nihai savunusunu nasıl yeniden kurduğu da değerlendirilmelidir. Süreç görünür hâle geldiğinde akademik dürüstlük bir tespit yazılımı sorunu olmaktan çıkar ve öğrencinin gerekçelendirme yükümlülüğünün parçasına dönüşür.

Tek ders kitabı modeli de dinamik ve uzman denetimli sayısal ders seçkilerine dönüşmelidir. Bir felsefe ders seçkisi giriş açıklaması, birincil metin, güncel akademik tartışma, kavram sözlüğü, argüman alıştırması, geri çağırma soruları ve aralıklı tekrar bölümleri içermelidir. Kerimoğlu’nun çalışması, büyük dil modellerinin insan dil ediniminin doğrudan kopyası olmadığını, fakat istatistiksel öğrenmenin hangi dilsel yapıları edinmeye yeterli olduğunu sınamak için güçlü bir yöntemsel laboratuvar sunduğunu göstermektedir (Kerimoğlu, 2026). Aynı ihtiyat Türkçe felsefe materyalleri için de geçerlidir: Model açıklama üretir; uzman doğrular, seçer ve bağlama yerleştirir.

Denetimsiz uzun yazılı ödev kutsal bir ölçme aracı değildir. Büyük dil modelleri çağında tek başına öğrencinin kavrayışını ölçmez; çoğu zaman öğrencinin araca erişimini, istem kurma becerisini ve öğretim üyesinin şüphe düzeyini ölçer. Yapay zekâ tespit yazılımını pedagojinin yerine koymak sayısal falcılıktır; nitekim Turnitin ve Originality gibi ticari dedektörleri 192 metinlik dengeli bir veri setiyle sınavan 2026 tarihli küçük ölçekli bir değerlendirme, bu araçların karma insan-yapay zekâ yazılarını ayırt etmede zorlandığına, doğruluğun metin türüne ve uzunluğuna duyarlı olduğuna ve İngilizceyi yabancı dil olarak yazan öğrencilerin metinleri bakımından adalet kaygılarının sürdüğüne işaret etmektedir (Hadra vd., 2026). Yazılı ödev korunacaksa aşamalı taslak, kaynak doğrulama, değişiklik günlüğü, yapay zekâ kullanım beyanı ve kısa sözlü savunmayla desteklenmelidir.

“Mill’in özgürlük anlayışını anlatan üç bin kelimelik bir metin yaz” görevi yerine şu görev verilebilir: “Büyük dil modelinin Mill’in zarar ilkesine ilişkin yorumundaki üç yanlış bul; özgürlük ile toplumsal zarar arasındaki ayrımı yeniden kur; güçlü bir paternalizm itirazı geliştir ve cevabını beş dakika sözlü olarak savun.” Böyle bir değerlendirme yazının cilasını değil, düşüncenin dayanıklılığını sınar.

Lisansüstü Eğitim ve Danışmanlığın Yeniden Kurulması

Yüksek lisans tezleri çoğu zaman “X düşünüründe Y kavramı” başlıklı genişletilmiş literatür özetlerine dönüşmektedir. Büyük dil modelleri bu tür metinleri üretmekte son derece başarılıdır; kötü tez modelini sürdürürsek yalnızca betimleyici metin enflasyonunu hızlandırırız. “William James’te bilinç” yerine “William James’in bilinç akışı görüşü çağdaş benlik kuramlarının süreklilik sorununa nasıl cevap verebilir?”, “Kuhn’da bilim” yerine “Kuhn’un ölçülemezlik savı bilimsel ilerleme kavramını bütünüyle terk etmeyi gerektirir mi?”, “Rorty’de hakikat” yerine “Rorty’nin temsilcilik eleştirisi epistemik normları açıklamakta yetersiz midir?” gibi sorular seçilmelidir. Tez açık bir problem, savunulabilir bir iddia, ciddi bir itiraz ve belirlenebilir bir katkı içermelidir.

Zwart’ın akademik yazarlık tartışması burada doğrudan önem taşır. Akademik yazarlık hiçbir zaman teknolojiden bağımsız olmamıştır; büyük dil modelleri yazarlığın anonimleşmesi ve kişisizleşmesi yönündeki daha eski eğilimleri hızlandırmaktadır (Zwart, 2026). Yazarlık, her kelimeyi çıplak elle yazmak değil; iddiayı anlayabilmek, kaynakları doğrulamak, metni savunmak, eleştiri karşısında yeniden kurmak ve sonuçların sorumluluğunu taşımaktır. Bu nedenle her tezde araştırma sorusu, literatür matrisi, argüman haritası, aşamalı taslaklar, yapay zekâ kullanım günlüğü ve sözlü savunma bulunmalıdır.

Doktora tezi büyük bir yüksek lisans tezi değil, uluslararası literatürde konumlanan bir araştırma programıdır. Her doktora öğrencisi tez metninin yanında en az bir makale taslağı, literatür matrisi, argüman haritası ve yapay zekâ kullanım kaydı üretmelidir. Danışmanın rolü, cümle düzeltmenin ötesine geçerek araştırma sürecinin mimarisini kurmaya yönelmelidir. Öğrenci toplantı öncesinde büyük dil modeliyle tezini savunup zayıf noktalarını saptayabilir; danışman toplantısı böylece metin düzeltme seansından gerçek felsefi denetime geçebilir.

Reform danışmanlıkta değil, öğrenci alımında başlamalıdır. Yüksek lisans ve doktora yazılı sınavları ezberlenmiş düşünür listelerini değil, canlı metin çözümleme, argüman yeniden kurma, kaynak güvenilirliği ve yeni probleme aktarım becerisini ölçmelidir. Adaya daha önce görmediği kısa bir Quine metni verilip temel iddiayı, çıkarım yapısını ve olası itirazı kurması istenebilir. Peirce’ın bir metni üzerinden kaçırma çıkarımının yapısı çözümlenebilir. Dennett’in yönelimsel tutumuna ilişkin kısa bir savunuyu verilerek araçsalcılık ile gerçekçilik arasındaki gerilim gösterilebilir.

Sözlü sınavlarda adaydan bir görüşü açıklaması, beklenmedik bir itiraza cevap vermesi, hata gösterildiğinde pozisyonunu gözden geçirmesi ve akademik İngilizce bir özet eleştirel biçimde değerlendirmesi istenmelidir. Sınavın bir bölümü yapay zekâ kullanımına kapalı, bir bölümü açıkça yapay zekâ destekli tasarlanabilir; geleceğin araştırmacısı hem bağımsız düşünmeli hem de araçla düşünmeyi bilmelidir.

Doktora yeterlik sınavı da ansiklopedik hafıza yarışından çıkarılmalıdır. Aday görülmemiş bir metni çözümlemeli, alanındaki temel tartışmaları ilişkilendirmeli, kendi araştırma programını sözlü olarak savunmalı ve büyük dil modeli tarafından üretilmiş sorunlu bir literatür özetini kaynak düzeyinde denetlemelidir. Tez İzleme Komitesi toplantıları yalnızca biçimsel ilerleme sunumlarıyla sınırlı kalmamalıdır. Her toplantıda güncel araştırma sorusu, argüman haritası, yeni literatür matrisi, başarısız denemeler, yapay zekâ kullanım günlüğü ve bir sonraki dönemin sınanabilir hedefleri incelenmelidir. Komite öğrencinin kaç sayfa yazdığını değil, araştırma problemini daha iyi anlayıp anlamadığını izlemelidir. Sayfa sayısı epistemik ilerleme birimi değildir.

Türkiye İçin Sıçrama İmkânı ve Kurumsal Riskler

Büyük dil modelleri akademik merkez ile çevre arasındaki bazı maliyetleri hızla düşürmektedir. Akademik İngilizce, metin düzenleme ve araştırma iş akışı maliyetlerinin düşmesinin, ana dili İngilizce olan ülkeler ile olmayan ülkeler arasındaki uluslararası yayın farkını daraltabileceği öngörülebilir. Bu, makalede kanıt olarak değil, bibliyometrik yöntemlerle sınanabilir ve sınanması gereken bir hipotez olarak ileri sürülmektedir.

Türkiye’de felsefe için mesele yalnızca daha düzgün İngilizce yazmak değildir. Öğrenci güncel literatüre daha hızlı girebilir, kavramsal harita çıkarabilir, metinler arası bağlantıları görebilir ve uluslararası hakem itirazlarına hazırlanabilir. Bölgesel bir üniversitedeki öğrenci, doğru kurumsal erişim sağlandığında, daha önce pahalı özel eğitim veya seçkin akademik ağlarla elde edilebilen desteğin bir kısmına ulaşabilir. Bu nedenle büyük dil modeli okuryazarlığı bireysel bir merak değil, epistemik altyapıdır.

Fakat bu araçların ticari platformlar içinde çalıştığı unutulmamalıdır. Nizam Özoğur, Börekçi ve Şeker, insan benzerliği ve bilinç söylemlerinin platform kapitalizmi, veri tekelleşmesi ve dikkat ekonomisinden bağımsız olmadığını vurgular (Nizam Özoğur vd., 2026). Bu eleştiri teknolojinin bütün kapasitesini ticari çıkarlara

indirmek için kullanılmamalıdır; fakat üniversitelerin eğitim altyapısını birkaç şirketin kapalı sistemlerine düşünmeden teslim etmemesi gerektiğini gösterir. Üniversiteler felsefeye özgü yapay zekâ pedagojisi, güvenli kurumsal modeller, doğrulanmış Türkçe ders seçkileri, açık kullanım standartları ve öğretim elemanı eğitimi geliştirmelidir. Aksi hâlde aynı eşitsizlik bu kez “iyi modeli, iyi istemi ve iyi İngilizceyi kim kullanabiliyor?” biçiminde geri döner.

Öğretim elemanlarının eğitimi bu kurumsal programın merkezinde yer almalıdır. Bir felsefe öğretim üyesinin her modelin teknik ayrıntısını bilmesi gerekmez; ancak modelin güçlü ve zayıf yönlerini deneyerek tanınması, öğrencinin hangi yollarla görev devrettiğini anlaması, doğrulama gerektiren çıktıları ayırt etmesi ve dersini buna göre yeniden tasarlaması gerekir. Bu yeterlilik kişisel merakı bırakıldığında bölümler arasındaki mevcut kalite farkları daha da büyüyebilir. Kurumsal lisanslar, disipline özgü örnek görev bankaları, düzenli uygulama atölyeleri ve açık etik ilkeler bu nedenle lüks değil, yükseköğretim altyapısının yeni parçalarıdır.

Bu önerilerin uygulanabilirlik düzeyleri aynı değildir. İzlenice tasarımı, görev bazlı yapay zekâ politikası ve ders içi değerlendirme biçimleri büyük ölçüde öğretim üyesi ve bölüm inisiyatifiyle değiştirilebilir. Lisansüstü kabul ölçütleri, yeterlik sınavının biçimi ve Tez İzleme Komitesi işleyişi ise enstitü yönergeleri ile Lisansüstü Eğitim ve Öğretim Yönetmeliği’nin (Yükseköğretim Kurulu, 2016) çizdiği çerçeve içinde kalmak zorundadır; bu düzeydeki değişiklikler senato kararı, bazı durumlarda mevzuat güncellemesi gerektirir. Kurumsal lisans önerisi de erişim eşitsizliği uyarısıyla birlikte okunmalıdır: Pahalı kapalı sistemlere erişemeyen kurumlar için açık ağırlıklı modeller ve kademeli erişim düzenlemeleri, teknolojinin eşitleyici vaadinin yeni bir eşitsizlik katmanına dönüşmemesinin zorunlu tamamlayıcılarıdır.

Sonuç

Büyük dil modelleri Türkiye’de felsefe eğitiminin krizini yaratmadı. Tek yönlü öğretimin, eski kaynakların, denetimsiz yazılı ödevlerin, betimleyici tezlerin ve biçimsel danışmanlığın epistemik yetersizliğini açığa çıkardı. Şimdi iki seçenek vardır: Eski sistemi yapay zeka tespit yazılımları ve etik belgelerle ayakta tutmaya çalışmak ya da nöropedagojik bulgular ışığında eğitimi yeniden kurmak.

Bu yeniden kuruluşun ilkesi açıktır. Büyük dil modeli öğrencinin yerine düşünen bir vekil değil, öğrenciyi daha fazla hatırlamaya, ayırmaya, doğrulamaya, itiraz etmeye ve savunmaya zorlayan bilişsel bir iskele olmalıdır. İzleniceden ders kitabına, lisansüstü öğrenci kabulünden doktora yeterliğine, Tez İzleme Komitesinden danışmanlığa kadar bütün yapı bu ilkeye göre yenilenmelidir.

Felsefe eğitiminin amacı düzgün görünen metin üretmek değil, hakikat karşısında sorumluluk taşıyan düşünür yetiştirmektir. Büyük dil modelleri bunu bizim yerimize yapmayacaktır. Fakat felsefe eğitimi dönüştürmek için güçlü bir araç sunabilir ve bu dönüşümü daha fazla ertelemeyi zorlaştırabilir.

Kaynakça

- Arf, C. (1959). Makine düşünebilir mi ve nasıl düşünebilir? *Atatürk Üniversitesi Halk Konferansları I*, 91-103.
- Bastani, H., Bastani, O., Sungu, A., Ge, H., Kabakcı, Ö. ve Mariman, R. (2025). Generative AI without guardrails can harm learning: Evidence from high school mathematics. *Proceedings of the National Academy of Sciences*, 122(26), e2422633122. <https://doi.org/10.1073/pnas.2422633122>
- Cepeda, N. J., Pashler, H., Vul, E., Wixted, J. T. ve Rohrer, D. (2006). Distributed practice in verbal recall tasks: A review and quantitative synthesis. *Psychological Bulletin*, 132(3), 354-380. <https://doi.org/10.1037/0033-2909.132.3.354>
- Dunlosky, J., Rawson, K. A., Marsh, E. J., Nathan, M. J. ve Willingham, D. T. (2013). Improving students' learning with effective learning techniques: Promising directions from cognitive and educational psychology. *Psychological Science in the Public Interest*, 14(1), 4-58. <https://doi.org/10.1177/1529100612453266>
- Erciyes, F. (2024). Turing testinin davranışçı ve işlevselci yorumu. *MetaZihin*, 7(1), 43-57. <https://doi.org/10.51404/metazihin.1280648>
- Fan, Y., Tang, L., Le, H., Shen, K., Tan, S., Zhao, Y., Shen, Y., Li, X. ve Gašević, D. (2025). Beware of metacognitive laziness: Effects of generative artificial intelligence on learning motivation, processes, and performance. *British Journal of Educational Technology*, 56(2), 489-530. <https://doi.org/10.1111/bjet.13544>
- Freeman, S., Eddy, S. L., McDonough, M., Smith, M. K., Okoroafor, N., Jordt, H. ve Wenderoth, M. P. (2014). Active learning increases student performance in science, engineering, and mathematics. *Proceedings of the National Academy of Sciences*, 111(23), 8410-8415. <https://doi.org/10.1073/pnas.1319030111>
- Gerlich, M. (2025). AI tools in society: Impacts on cognitive offloading and the future of critical thinking. *Societies*, 15(1), Article 6 (Düzeltilme: *Societies*, 15(9), Article 252). <https://doi.org/10.3390/soc15010006>
- Güven, Ö. (2026). Anlayan bir yapay zekâ olanaklı mı? Ö. Güven ve F. D. Çoban Sarı (Ed.), *Yapay zekâyâ felsefi yaklaşımlar* içinde (ss. 25-40). İstanbul: İstanbul Üniversitesi Yayınevi. <https://doi.org/10.26650/BS/AH11T3.2026.002-5.03>
- Güven, Ö. ve Çoban Sarı, F. D. (Ed.). (2026). *Yapay zekâyâ felsefi yaklaşımlar*. İstanbul: İstanbul Üniversitesi Yayınevi. <https://doi.org/10.26650/BS/AH11T3.2026.002-5>
- Hadra, M., Cambridge, K. ve Mesbah, M. (2026). Evaluating the accuracy and reliability of AI content detectors in academic contexts. *International Journal for Educational Integrity*, 22, Article 4. <https://doi.org/10.1007/s40979-026-00213-1>
- Hicks, M. T., Humphries, J. ve Slater, J. (2024). ChatGPT is bullshit. *Ethics and Information Technology*, 26, 38. <https://doi.org/10.1007/s10676-024-09775-5>
- Kerimoğlu, C. (2026). Yapay zekâ, dilin doğuştanlığını yıkabilir mi? Büyük dil modelleri karşısında uyarıcı eksikliği argümanı. *Dil Araştırmaları*, 38, 1-24. <https://doi.org/10.54316/dilarastirmalari.1879242>
- Kozanitis, A. ve Nenciovici, L. (2023). Effect of active learning versus traditional lecturing on the learning achievement of college students in humanities and social sciences: A meta-analysis. *Higher Education*, 86, 1377-1394. <https://doi.org/10.1007/s10734-022-00977-8>
- Kuşçu, E. S. (2026). Doğal zekânın yapaylaştırılması üzerine yeni bir kötümser yaklaşım. Ö. Güven ve F. D. Çoban Sarı (Ed.), *Yapay zekâyâ felsefi yaklaşımlar* içinde (ss. 93-106). İstanbul: İstanbul Üniversitesi Yayınevi. <https://doi.org/10.26650/BS/AH11T3.2026.002-5.07>
- Lake, B. M., Ullman, T. D., Tenenbaum, J. B. ve Gershman, S. J. (2017). Building machines that learn and think like people. *Behavioral and Brain Sciences*, 40, e253. <https://doi.org/10.1017/S0140525X16001837>
- Millière, R. ve Buckner, C. (2026). The philosophy of language models. *Philosophy Compass*, 21, e70095. <https://doi.org/10.1111/phc3.70095>

- Monett, D., Lewis, C. W. P. ve Thórisson, K. R. (2020). Introduction to the JAGI special issue “On defining artificial intelligence”. *Journal of Artificial General Intelligence*, 11(2), 1-4. <https://doi.org/10.2478/jagi-2020-0003>
- Nizam Özoğur, H., Börekçi, F. H. ve Şeker, Ş. E. (2026). Yapay zekâ ve bilinç: Felsefi perspektifler. Ö. Güven ve F. D. Çoban Sarı (Ed.), *Yapay zekâyâ felsefi yaklaşımlar* içinde (ss. 75-92). İstanbul: İstanbul Üniversitesi Yayınevi. <https://doi.org/10.26650/BS/AH11T3.2026.002-5.06>
- Oakley, B., Rogowsky, B. ve Sejnowski, T. J. (2021). *Uncommon sense teaching: Practical insights in brain science to help students learn*. New York: TarcherPerigee.
- Rasch, B. ve Born, J. (2013). About sleep's role in memory. *Physiological Reviews*, 93(2), 681-766. <https://doi.org/10.1152/physrev.00032.2012>
- Roediger, H. L. ve Karpicke, J. D. (2006). Test-enhanced learning: Taking memory tests improves long-term retention. *Psychological Science*, 17(3), 249-255. <https://doi.org/10.1111/j.1467-9280.2006.01693.x>
- Sarı, Ö. O. (2026). Semantik içerik ne zaman ortaya çıkar? Çin Odası, büyük dil modelleri ve nedensel risk. Ö. Güven ve F. D. Çoban Sarı (Ed.), *Yapay zekâyâ felsefi yaklaşımlar* içinde (ss. 41-56). İstanbul: İstanbul Üniversitesi Yayınevi. <https://doi.org/10.26650/BS/AH11T3.2026.002-5.04>
- Suriano, R., Plebe, A., Acciai, A. ve Fabio, R. A. (2025). Student interaction with ChatGPT can promote complex critical thinking skills. *Learning and Instruction*, 95, Article 102011. <https://doi.org/10.1016/j.learninstruc.2024.102011>
- Turing, A. M. (1948). *Intelligent machinery*. London: National Physical Laboratory. Yeniden basım: B. J. Copeland (Ed.). (2004). *The essential Turing* içinde (ss. 395-432). Oxford: Oxford University Press.
- Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., Kaiser, Ł. ve Polosukhin, I. (2017). Attention is all you need. *Advances in Neural Information Processing Systems*, 30, 5998-6008.
- Yükseköğretim Kurulu. (2016). Lisansüstü Eğitim ve Öğretim Yönetmeliği. Resmî Gazete, 20/4/2016, Sayı: 29690 (En son değişiklik: 6/3/2026).
- Zwart, H. (2026). AI and the death of the academic author: A dialectical assessment of the challenges of generative AI for the concept and practice of academic authorship. Ö. Güven ve F. D. Çoban Sarı (Ed.), *Yapay zekâyâ felsefi yaklaşımlar* içinde (ss. 1-11). İstanbul: İstanbul Üniversitesi Yayınevi. <https://doi.org/10.26650/BS/AH11T3.2026.002-5.01>

Extended Abstract

Purpose

This article examines how large language models can be used to redesign philosophy education in Türkiye rather than merely policed as instruments of academic dishonesty. The central claim is that the arrival of large language models has not created the deepest crisis in philosophy education; it has exposed long-standing weaknesses that were already present before the current technological rupture. These weaknesses include one-way lecturing, outdated syllabi, limited engagement with current international scholarship, descriptive graduate theses, irregular supervision, and assessment practices that treat a polished written product as sufficient evidence of philosophical understanding. The article therefore asks a practical and normative question: under what conditions can large language models become tools for better philosophical learning instead of shortcuts that weaken epistemic agency?

The argument is motivated by the contrast between two common but inadequate responses to generative artificial intelligence. The first response treats the technology primarily as a cheating machine and reduces institutional policy to detection, prohibition, and disciplinary anxiety. The second response overstates the cognitive status of current systems and assumes that fluent text production is already equivalent to human-like understanding. Against both views, the article proposes a middle position. Large language models should neither be mythologised as autonomous epistemic subjects nor dismissed as useless statistical parrots. They should be understood as powerful but fallible cognitive tools whose pedagogical value depends on how they are embedded in learning design, assessment, supervision, and institutional policy.

Methodology

The study is a conceptual and normative analysis. It does not report original empirical classroom data; instead, it synthesises work in philosophy of artificial intelligence, cognitive psychology, neuropedagogy, and higher education pedagogy in order to formulate a reform framework for philosophy education. The article first clarifies the epistemic status of large language models by distinguishing artificial intelligence in general from large language model-based systems in particular. It then connects this clarification to philosophical debates about understanding, semantic content, authorship, and human-like learning. The article draws on Turing's early idea of a trainable "child machine" and Arf's reflections on memory, understanding, and adaptive intelligence to show that the connection between machine intelligence and education is not historically new (Turing, 1948; Arf, 1959).

The pedagogical framework is drawn mainly from research on learning, memory, and active engagement. Oakley, Rogowsky, and Sejnowski's neuropedagogical account is used to emphasise that learning is not the mere feeling of following an explanation, but the organisation, consolidation, retrieval, and transfer of knowledge in long-term memory (Oakley et al., 2021). This point is strengthened by evidence that retrieval practice and spaced learning are among the most effective learning techniques, whereas rereading and highlighting have more limited value when used alone (Dunlosky et al., 2013). The article also relies on meta-analytic evidence that active learning improves student performance, including evidence concerning humanities and social sciences contexts (Freeman et al., 2014; Kozanitis & Nenciovici, 2023).

Philosophically, the article treats model fluency and model understanding as distinct issues. It uses work on the philosophy of language models to reject the simplistic view that the training objective of next-token prediction fully explains all post-training capacities (Millière & Buckner, 2026). At the same time, it resists the opposite inference that successful text generation proves robust understanding, belief, consciousness, or semantic responsibility. Discussions of understanding, causal risk, bullshit, authorship, and platform capitalism are used to frame large language models as tools that require external verification and normative human responsibility (Güven, 2026; Hicks et al., 2024; Nizam Özoğur et al., 2026; Sarı, 2026; Zwart, 2026).

Findings

The first finding is that philosophy education cannot be protected by merely banning or detecting artificial intelligence use. Take-home essays, isolated written assignments, and descriptive literature summaries are no longer reliable proxies for philosophical competence. More importantly, they were already weak proxies before

large language models became widely available. A student can produce a fluent essay on Hume, Mill, Quine, Kuhn, Dennett, or Rorty without being able to reconstruct an argument, identify a hidden premise, formulate a strong objection, or defend a thesis under pressure. Large language models intensify this problem because they can generate polished descriptions rapidly; but the underlying pedagogical failure is the reduction of philosophy to information storage rather than argumentative practice.

The second finding is that large language models can be pedagogically valuable when they are used to strengthen memory, retrieval, argumentation, and verification. They should not be used as answer-generating machines that replace student effort. Instead, they can support personalised retrieval practice, spaced repetition, interleaved questioning, graduated hints, argument rehearsal, source checking, and oral defence preparation. For example, instead of asking a model to write an explanation of Hume's critique of necessary connection, a student can ask it to test recall without immediately revealing the answer, provide hints after errors, return to the same question after several days, and combine it with Mill's account of induction. Such uses turn the model from a shortcut into a cognitive scaffold.

The third finding concerns course design. A syllabus should not be a bureaucratic list of authors and topics; it should be an epistemic architecture. Each week should specify the philosophical problem, primary text, concepts to be retrieved, arguments to be reconstructed, permissible or required uses of large language models, and forms of live performance through which understanding will be tested. In the philosophy of science, for instance, a course should not simply list Quine, Kuhn, and Rorty. It should organise the weeks around problems such as theory-ladenness of observation, scientific change, incommensurability, truth, and justification. Students should learn to compare positions, reconstruct disputes, and defend interpretations, not merely report what canonical figures said.

The fourth finding is institutional. Assessment should combine artificial-intelligence-closed and artificial-intelligence-open tasks. Closed tasks, such as in-class writing, live text analysis, and oral examination, test independent recall and philosophical reasoning. Open tasks test the student's ability to use models responsibly: formulating questions, checking sources, detecting false claims, rejecting weak suggestions, documenting prompts, and revising arguments. In this design, academic honesty is not outsourced to unreliable detection software. It becomes part of the student's burden of justification.

The fifth finding concerns graduate education. Master's and doctoral research should move away from descriptive "X philosopher on Y concept" projects and toward problem-driven work that contains a clear question, a defensible thesis, serious objections, and an identifiable contribution. Large language models make descriptive thesis inflation easier, so supervision must shift from sentence-level correction to the architecture of inquiry: research questions, literature matrices, argument maps, staged drafts, model-use logs, and oral defences. Admissions, qualifying exams, and thesis monitoring committees should likewise test live analysis, argument reconstruction, source reliability, and the ability to respond to unexpected objections.

Limitations

The article is limited in four respects. First, it is conceptual rather than empirical. Its proposals should be tested through classroom interventions, comparative course designs, student interviews, and assessment studies. Second, the discussion focuses on philosophy education in Türkiye, although many of the problems it identifies are shared by other humanities and social sciences contexts. Third, the article does not settle metaphysical questions about machine consciousness, genuine understanding, or subjective experience. These questions are important, but they are not identical with the pedagogical question of whether a system can help students retrieve, compare, criticise, and justify philosophical claims. Fourth, the article assumes that access, infrastructure, and institutional policy matter. If only privileged students have reliable access to advanced models, prompt literacy, academic English support, and high-quality feedback, the same technology that could reduce epistemic inequality may reproduce it in a new form.

Implications

The practical implication is that universities should stop treating large language models as an external contamination of an otherwise healthy system. Philosophy departments need explicit task-level artificial intelligence policies, model-use declarations, revised assessment designs, and training for instructors. Some tasks

should prohibit model use; some should allow it in limited forms; some should require it in order to teach verification and critique. Departments should develop discipline-specific prompt examples, validated Turkish course materials, retrieval-practice routines, and safe institutional workflows. This is not a matter of technological enthusiasm. It is a matter of aligning philosophy education with what is known about learning, memory, feedback, and the contemporary research environment.

The article also implies that academic authorship should be understood as responsibility rather than mere manual production of every sentence. A scholar remains responsible for claims, sources, inferences, omissions, and final judgments even when digital tools are used. For students, this means that model-assisted work must be accompanied by source verification, revision records, and oral defence. For supervisors, it means that the main task is no longer only to correct wording but to evaluate whether the student understands the problem, can defend the argument, and can explain which tool suggestions were accepted or rejected and why.

Originality and Value

The originality of the article lies in connecting three discussions that are often kept apart: the philosophy of large language models, neuropsychological findings about learning and memory, and the concrete institutional weaknesses of philosophy education in Türkiye. Rather than asking only whether machines understand or whether students cheat, the article reframes large language models as a test of the epistemic design of philosophy education itself. Its value is therefore both theoretical and practical. Theoretically, it clarifies why model fluency does not remove the need for human justification and why the absence of human-like consciousness does not make the model pedagogically irrelevant. Practically, it offers a reform agenda for syllabi, assignments, graduate admissions, doctoral qualification exams, thesis monitoring committees, and supervision. The central conclusion is that large language models will not think for philosophy education. Used properly, however, they can force philosophy education to become more demanding, more dialogical, more memory-sensitive, and more epistemically responsible.