

BIGS AI GOVERNANCE FRAMEWORK

Kurumsal BT Operasyonlarında Yapay Zekâ Otonomisini
Güvenle Yönetmek İçin Bir Karar Çerçevesi

Whitepaper • Sürüm 1.0 • Temmuz 2026

Bigs Bilişim Limited Şirketi

www.bigsbilisim.com

1. Yönetici Özeti

Yapay zekâ ajanları (AI agents) kurumsal BT operasyonlarına hızla giriyor. Gartner, 2026 sonunda kurumsal uygulamaların %40'ının göreve özel AI ajanları içereceğini öngörüyor; bu oran 2025'te yalnızca %5'ti. Ancak çoğu kurumun şu temel soruya ölçülebilir bir cevabı yok: Bir AI ajanı hangi işlemi kendi başına yapabilir, hangisinde insan onayı gerekir?

Bu sorunun cevapsız kalması iki uç noktadan birine savrulmaya yol açar: ya kontrolsüz otonomi (AI'nın üretim ortamında öngörülemeyen müdahaleler yapması) ya da sıfır otonomi (AI'nın pahalı bir gösterge panosuna indirgenmesi ve yatırımın karşılığının alınamaması).

Bigs AI Governance Framework, bu ikilemi çözmek için geliştirilmiş, üretim ortamlarında test edilmiş bir karar çerçevesidir. Çerçevenin kalbinde **Runbook Güven Endeksi (RCI)** yer alır: her operasyonel aksiyonu beş boyutta puanlayan ve dört net otonomi seviyesinden birine haritalayan ölçülebilir bir güven modeli. Çerçeve ayrıca operasyonel güvenlik mekanizmaları (15 Dakika Kurtarma Kuralı, Değişiklik Risk Puanlama, denetim izi), bir Operasyonel Bellek Katmanı ve otonominin zamanla kazanıldığı bir runbook yaşam döngüsü tanımlar.

Bu doküman, çerçevenin ilkelerini, puanlama modelini, güvenlik mekanizmalarını ve teknoloji-bağımsız referans mimarisini sunar; ardından kurumların çerçeveyi kendi ortamlarına uygulaması için dört fazlı bir yol haritası önerir.

2. Problem: Kontrolsüz Otonomi ile Sıfır Otonomi Arasında

BT operasyon ekipleri yıllardır aynı döngüde: alarm gelir, mühendis uyanır, teşhis koyar, çoğu zaman daha önce onlarca kez uygulanmış standart bir müdahaleyi tekrar eder. Sektör araştırmaları, BT triyaj işlemlerinin ve tekrarlayan müdahalelerin %40-60'ının 2026 itibarıyla otomatikleştirilebilir olduğunu gösteriyor. Büyük dil modelleri ve ajan mimarileri, bu otomasyonu teknik olarak mümkün kılıyor.

Teknik olarak mümkün olan ile operasyonel olarak güvenli olan arasındaki fark ise yönetişimdir. Sahada iki tipik başarısızlık modeli gözlemliyoruz:

- **Kontrolsüz otonomi:** AI'a geniş yetkiler verilir; sistem, etki alanı büyük bir işlemi (örneğin bir veritabanı failover'ı) yanlış bir alarmla dayanarak gece 03:00'te tetikler. Tek bir kötü deneyim, kurumun AI otomasyonuna güvenini yıllarca geriye atar.
- **Sıfır otonomi:** Risk korkusuyla AI yalnızca 'öneri' üretir, her adım insana sorulur. Mühendis iş yükü azalmaz, yanıt süreleri kısalmaz; AI yatırımı gösterişli bir sohbet arayüzüne dönüşür.

Her iki mod da aynı eksiklikten kaynaklanır: **ölçülebilir bir güven modeli** yoktur. 'Bu işleme ne kadar güveniyoruz?' sorusu sezgiyle değil, puanla cevaplanmalıdır. Bigs AI Governance Framework tam olarak bunu yapar.

3. Temel İlkeler

Çerçeve beş ilke üzerine kuruludur. Tüm mekanizmalar bu ilkelerin uygulamasıdır:

- **İlke 1 — Otonomi kazanılır, varsayılmaz.** Hiçbir aksiyon 'muhtemelen güvenlidir' diye otonom çalıştırılmaz. Otonomi, ölçülebilir kriterlerin ve tarihsel kanıtın sonucudur.
- **İlke 2 — Her otonom aksiyon geri alınabilir olmalıdır.** Geri dönüşü olmayan veya geri dönüşü belirsiz işlemler, başarı geçmişi ne olursa olsun tam otonomiye terfi edemez.

- **İlke 3 — Belleği olmayan AI'a güvenilmez.** AI, her olayı sıfırdan değerlendirmemeli; geçmiş olayları, kök nedenleri ve müdahale sonuçlarını bağlam olarak almalıdır.
- **İlke 4 — İnsan her zaman devrededir.** Otonomi, insanın devre dışı kalması değil, insanın doğru anda ve doğru bilgiyle devreye girmesidir. Eskalasyon yolları her aksiyon için önceden tanımlıdır.
- **İlke 5 — Her karar izlenebilirdir.** AI'ın ne yaptığı kadar neden yaptığı da kayıt altındadır: tetikleyici, skor, karar, sonuç ve süre.

4. Runbook Güven Endeksi (RCI)

RCI, bir operasyonel aksiyonun (runbook'un) AI tarafından hangi otonomi seviyesinde çalıştırılabileceğini belirleyen puanlama modelidir. Her runbook beş boyutta 0-20 arası puanlanır; toplam skor 0-100 aralığındadır.

4.1 Beş Boyut

Boyut	Sorduğu Soru	Yüksek Puan Örneği
Öngörülebilirlik	Aksiyonun sonucu ne kadar deterministik? Aynı girdi her zaman aynı çıktıyı üretir mi?	Servis yeniden başlatma
Geri Alınabilirlik	İşlem yanlış giderse ne kadar hızlı ve tam olarak geri alınabilir?	Konfigürasyon yedeği alınan değişiklik
Etki Alanı	İşlem tek bir servisi mi, sunucuyu mu, yoksa tüm ortamı mı etkiler?	Tek uygulama havuzu (dar kapsam)
Tarihsel Başarı	Bu runbook geçmişte kaç kez çalıştı ve başarı oranı nedir?	Son 50 çalıştırmada %100 başarı
Tespit Güveni	Aksiyonu tetikleyen alarm ne kadar güvenilir? Yanlış pozitif oranı düşük mü?	Çok kaynaklı doğrulanmış alarm

4.2 Dört Otonomi Seviyesi

Toplam RCI skoru, aksiyonu dört karar seviyesinden birine haritalar:

RCI Skoru	Seviye	Davranış
85 – 100	Otonom	AI aksiyonu kendi başına çalıştırır, sonucu raporlar. İnsan yalnızca bilgilendirilir.
65 – 84	AI Önerisi	AI müdahaleyi hazırlar ve gerekçelendirir; insan tek onayla çalıştırır.
40 – 64	Onay Gerekli	İnsan, AI'nı teşhisini ve önerisini ayrıntılı inceler; gerekçeli onay veya ret verir.
0 – 39	Yalnızca Manuel	AI hiçbir aksiyon almaz; yalnızca teşhis, bağlam ve geçmiş olay bilgisi sunar.

4.3 Örnek Senaryolar

Aksiyon	RCI	Seviye	Belirleyici Etken
IIS uygulama havuzu yeniden başlatma	87	Otonom	Dar etki, yüksek tarihsel başarı
Geçici dosya temizliği ile disk alanı açma	78	AI Önerisi	Öngörülebilir, ancak silme geri alınamaz
Üretim uygulama sunucusunda Windows servisi yeniden başlatma	58	Onay Gerekli	Etki alanı geniş, bağımlılıklar var

Aksiyon	RCI	Seviye	Belirleyici Etken
Üretim veritabanı failover işlemi	34	Yalnızca Manuel	Kritik etki, geri dönüş maliyeti yüksek

Puanlama, kurulum aşamasında BT ekibi ile birlikte yapılan bir çalıştayda belirlenir ve yaşam döngüsü boyunca (Bölüm 7) veriye dayalı olarak güncellenir. RCI'nin değeri kesin matematiksel doğruluğunda değil, otonomi tartışmasını sezgiden ölçüme taşımasındadır.

5. Operasyonel Güvenlik Mekanizmaları

5.1 15 Dakika Kurtarma Kuralı

Otonom veya onaylı bir müdahale başladıktan sonra sistem 15 dakika içinde stabil hale gelmezse, AI müdahaleyi durdurur ve olayı tüm bağlamıyla (yapılan işlemler, gözlemlenen etkiler, geçerli hipotezler) insan mühendise devreder. Bu kural, 'AI'nın kendi hatasını düzeltmeye çalışırken durumu derinleştirmesi' riskini yapısal olarak ortadan kaldırır.

5.2 AI Destekli Değişiklik Risk Puanlama

Çerçeve yalnızca olay müdahalesini değil, planlı değişiklikleri de kapsar. Her değişiklik talebi; etkilenen sistemler, değişiklik penceresi, geri alma planının varlığı ve benzer geçmiş değişikliklerin sonuçları üzerinden AI tarafından risk puanına bağlanır. Yüksek riskli değişiklikler otomatik olarak ek onay katmanına yönlendirilir.

5.3 Yetki Sınırlama ve Kontrollü Uygulama Katmanı

AI, hedef sistemlere doğrudan ve sınırsız erişimle değil, kontrollü bir uygulama katmanı üzerinden erişir: önceden tanımlı komut setleri (allowlist), aksiyon başına ayrı kimlik ve yetki, ortam bazlı ayırım (üretim / test) ve zaman pencereli erişim. AI'nın 'yapabileceği' şeyler kümesi, 'yapması onaylanan' şeyler kümesiyle sınırlıdır.

5.4 Denetim İzi (Audit Trail)

Her aksiyon için şu kayıtlar tutulur: tetikleyen alarm ve kaynağı, hesaplanan RCI skoru ve boyut kırılımı, verilen karar (otonom / onaylı / reddedildi), çalıştırılan komutlar, gözlemlenen sonuç ve toplam çözüm süresi. Bu kayıtlar hem uyumluluk gereksinimlerini karşılar hem de Bölüm 7'deki yaşam döngüsünün veri kaynağıdır.

6. Operasyonel Bellek Katmanı

Büyük dil modellerinin kurumsal operasyonlardaki en büyük zafiyeti bağlamsızlıktır: model, ortamınızın geçmişini bilmeden her olayı ilk kez görüyormuş gibi değerlendirir. Operasyonel Bellek Katmanı bu zafiyeti kapatır.

Katman; olay geçmişini, kök neden analizlerini, uygulanan müdahaleleri ve sonuçlarını yapılandırılmış ve vektörel olarak saklar. Yeni bir olay geldiğinde AI, benzer geçmiş olayları anlamsal olarak bulur ve karar bağlamına dahil eder: 'Bu sunucuda benzer bellek alarmı son 6 ayda 4 kez görüldü; 3'ünde kök neden X servisiydi ve Y müdahalesi çözdü.'

Bellek katmanı aynı zamanda RCI'nin Tarihsel Başarı ve Tespit Güveni boyutlarını besleyen veri kaynağıdır; yani bellek olmadan güven modeli de statik kalır. Bu, İlke 3'ün doğrudan uygulamasıdır.

7. Runbook Yaşam Döngüsü: Otonomi Nasıl Kazanılır?

Çerçeveyi statik bir politika dokümanından yaşayan bir sisteme dönüştüren unsur, otonomi seviyelerinin zamanla ve veriyle değişebilmesidir:

- **Başlangıç:** Yeni tanımlanan her runbook, ilk RCI puanlamasına göre en fazla 'Onay Gerekli' seviyesinde başlar. Hiçbir runbook doğrudan 'Otonom' doğmaz.
- **Terfi:** Her başarılı çalıştırma Tarihsel Başarı puanını artırır. Tanımlı eşikler aşıldığında sistem bir terfi önerisi üretir; terfi kararını her zaman insan verir.
- **Seviye düşürme:** Başarısız bir çalıştırma veya 15 Dakika Kuralı ihlali, runbook'u otomatik olarak bir alt seviyeye düşürür ve kök neden analizi tamamlanana kadar orada tutar.
- **Periyodik gözden geçirme:** Tüm runbook portföyü üç ayda bir gözden geçirilir; ortam değişiklikleri (yeni bağımlılıklar, mimari değişimler) puanlara yansıtılır.

Bu döngü, kurum içinde AI otonomisine güvenin de kademeli olarak inşa edilmesini sağlar: ekip, otonominin keyfi değil kazanılmış olduğunu somut verilerle görür.

8. Referans Mimari

Çerçeve teknoloji-bağımsızdır; aşağıdaki katmanlar farklı ürünlerle gerçekleştirilebilir. Parantez içindeki örnekler, Bigs Bilişim'in üretim ortamlarında doğrulanmış tercihleridir:

Katman	Görevi	Örnek Teknoloji
1. İzleme ve Tespit	Metrik, log ve olay toplama; alarm üretimi	Zabbix, özel ajanlar
2. Orkestrasyon	Olay akışının yönetimi, entegrasyonlar, onay akışları	n8n
3. AI Karar Katmanı	Teşhis, RCI değerlendirmesi, müdahale planı üretimi	LLM (API veya yerel)
4. Operasyonel Bellek	Olay geçmişi, kök neden ve sonuçların vektörel saklanması	PostgreSQL + pgvector
5. Kontrollü Uygulama	Yetki sınırlı, denetlenebilir komut çalıştırma	SSH / PowerShell
6. Görünürlük	Gerçek zamanlı durum, denetim izi ve yönetim raporları	Grafana

Mimarinin kritik özelliği, tüm bileşenlerin müşteri kontrolündeki altyapıda (on-premise veya müşterinin kendi bulutu) çalışabilmesidir: operasyonel veriler ve sistem erişimi üçüncü taraf bir SaaS platformuna devredilmek zorunda değildir. Bu, veri egemenliği ve KVKK gereksinimleri olan kurumlar için belirleyici bir tasarım kararıdır.

9. Uygulama Yol Haritası

Çerçevenin bir kuruma uyarlanması dört fazda gerçekleşir:

- Faz 1 — Envanter ve Runbook Çıkarımı (2-3 hafta):** Mevcut alarm türleri, tekrarlayan olaylar ve fiilî müdahale pratikleri belgelenir; ilk runbook portföyü çıkarılır.
- Faz 2 — RCI Puanlama Çalıştayı (1 hafta):** BT ekibiyle birlikte her runbook beş boyutta puanlanır; otonomi seviyeleri ve eskalasyon yolları tanımlanır.
- Faz 3 — Gölge Mod (4-8 hafta):** AI tüm olaylarda teşhis ve müdahale önerisi üretir ancak hiçbir aksiyon almaz; önerileri insan uygular. Bu dönem, Tespit Güveni ve öneri kalitesinin ölçüldüğü kanıt toplama evresidir.
- Faz 4 — Kademeli Otonomi:** Gölge mod verisiyle doğrulanan runbook'lar, yaşam döngüsü kuralları çerçevesinde önce 'AI Önerisi'ne, ardından hak edenler 'Otonom' seviyesine terfi eder.

10. Sonuç

Önümüzdeki iki yılın ayrıştırıcı sorusu 'AI kullanıyor musunuz?' değil, 'AI'nıza hangi kanıtla, ne kadar yetki veriyorsunuz?' olacak. Bigs AI Governance Framework bu soruya ölçülebilir bir cevap verir: otonomi kazanılır, her aksiyon geri alınabilir ve izlenebilirdir, insan her zaman devrededir ve güven, sezgiyle değil veriyle inşa edilir.

Çerçeve, Bigs Bilişim'in kendi AI destekli izleme ve operasyon platformunun geliştirilmesi sırasında, gerçek üretim ortamlarının gereksinimlerinden damıtılmıştır ve yeni saha deneyimleriyle güncellenmeye devam edecektir.

Bigs Bilişim Hakkında

Bigs Bilişim, 26 yılı aşan kurumsal sistem mühendisliği deneyimini (Windows, Linux, bulut ve SAP altyapıları) modern otomasyon ve yapay zekâ teknolojileriyle birleştiren İstanbul merkezli bir teknoloji şirkettir. Kurumsal müşterilerine AI destekli izleme ve operasyon hizmetleri, altyapı danışmanlığı ve otomasyon çözümleri sunar.

İletişim ve daha fazla bilgi: www.bigsbilisim.com