

06.08.2025 – 13:00 Uhr

Myrtle.ai ermöglicht ML-Inferenz-Latenzen im Mikrosekundenbereich mit VOLLO auf Napatech SmartNICs

Cambridge, England (ots/PRNewswire) -

[Myrtle.ai](#), ein anerkannter Marktführer bei der Beschleunigung von maschinellen Lernverfahren, hat heute die Unterstützung seines VOLLO®-Inferenzbeschleunigers auf der NT400D1x-Serie von SmartNICs von Napatech bekanntgegeben.

VOLLO erreicht branchenführende Latenzzeiten für ML-Inferenzberechnungen, die weniger als eine Mikrosekunde betragen können. Diese neue Version ermöglicht es denjenigen, die geringstmögliche Latenzen benötigen, Inferenzen direkt am Netzwerkrand auf einer SmartNIC auszuführen. Auf VOLLO kann eine Vielzahl von Modellen ausgeführt werden, darunter LSTM, CNN, MLP sowie Random Forests und Gradient Boosting-Entscheidungsbäume.

Dies wurde entwickelt, um den Anforderungen einer breiten Palette von Anwendungen gerecht zu werden, darunter Finanzhandel, Mobilfunk, Cybersicherheit, Netzwerkmanagement und andere, bei denen ML-Inferenzen mit der geringstmöglichen Latenzzeit ausgeführt werden, was Vorteile in Bezug auf Sicherheit, Profit, Effizienz und Kosten mit sich bringt.

„Wir freuen uns über die Zusammenarbeit mit dem Weltmarktführer im SmartNIC-Vertrieb, um beispiellos niedrige Latenzen für ML-Inferenzen zu ermöglichen“, sagt Peter Baldwin, CEO von Myrtle.ai. „Unsere Kunden sind aktiv auf der Suche nach immer niedrigeren Latenzen und diese neue Version wird es ihnen ermöglichen, den vollen Nutzen aus den führenden Latenzzeiten von VOLLO zu ziehen.“

Jarrold J.S. Siket, Chief Product & Marketing Officer bei Napatech, ist von den Aussichten begeistert. „Wir haben erkannt, dass die führenden Latenzzeiten in den STAC® ML-Benchmarks unseren Kunden im Finanzmarkt einen echten Mehrwert bieten können, da sie ML verstärkt für den automatischen Handel einsetzen“, sagt er. „Der VOLLO-Compiler soll es ML-Entwicklern sehr einfach machen, unsere SmartNICs zu nutzen, und das stärkt unser Produkt- und Dienstleistungsportfolio.“

Interessenten können den ML-orientierten VOLLO-Compiler ab sofort unter vollo.myrtle.ai herunterladen und herausfinden, welche Latenzen mit ihren Modellen auf den SmartNICs der NT400D1x-Serie von Napatech erreicht werden können.

Informationen zu Myrtle.ai

Myrtle.ai ist ein KI/ML-Softwareunternehmen, das erstklassige Inferenzbeschleuniger auf FPGA-basierten Plattformen aller führenden FPGA-Anbieter liefert. Myrtle verfügt über Fachwissen im Bereich neuronaler Netze für das gesamte Spektrum von ML-Netzen und hat Beschleuniger für FinTech, Sprachverarbeitung und Empfehlungen entwickelt.

„STAC“ und alle STAC-Namen sind Marken oder eingetragene Marken des Securities Technology Analysis Center, LLC.

Foto: https://mma.prnewswire.com/media/2743540/Myrtleai_VOLLO.jpg

Logo: https://mma.prnewswire.com/media/2739186/5443721/Myrtle_ai_Logo.jpg

View original content: <https://www.prnewswire.com/news-releases/myrtleai-ermoglicht-ml-inferenz-latenzen-im-mikrosekundenbereich-mit-vollo-auf-napatech-smartnics-302523075.html>

Pressekontakt:

Giles Peckham,
+44 7785 278478,
giles@myrtle.ai

Diese Meldung kann unter <https://www.presseportal.ch/de/pm/100102576/100933886> abgerufen werden.